

Hope, Fear and Aspirations*

Xue Dong He[†] and Xun Yu Zhou[‡]

July 23, 2012

Abstract

We propose a rank-dependent portfolio choice model in continuous time that captures the role in decision making of three emotions: hope, fear and aspirations. Hope and fear are modeled through a reversed S-shaped probability weighting function and aspirations through a probabilistic constraint. By employing the recently developed approach of quantile formulation, we solve the portfolio choice problem both thoroughly and analytically. These solutions motivate us to introduce a fear index, a hope index and a lottery-likeness index to quantify the impacts of three emotions, respectively, on investment behavior. We find that a sufficiently high level of fear endogenously necessitates portfolio insurance. On the other hand, hope is reflected in the agent's perspective on good states of the world: a higher level of hope causes the agent to include more scenarios under the notion of good states and leads to greater payoffs in sufficiently good states. Finally, an exceedingly high level of aspirations results in the construction of a lottery-type payoff, indicating that the agent needs to enter into a pure gamble in order to achieve his goal. We also conduct numerical experiments to demonstrate our findings.

Key words. portfolio choice, continuous time, rank-dependent utility, probability weighting, SP/A theory, quantile formulation, portfolio insurance

*We are grateful for comments from seminar and conference participants at the Chinese University of Hong Kong, Columbia University, Tongji University, the University of St. Gallen, Risk and Stochastics Day 2010 at the London School of Economics and Political Science, the 2010 Workshop on Stochastics, Control and Finance at Imperial College, the Sixth World Congress of the Bachelier Finance Society in Toronto, the Sixth Oxford-Princeton Workshop on Mathematical Finance at the University of Oxford, the 2011 SIAM Conference on Control and Its Applications in Baltimore, the 2011 Workshop of the Sino-French Summer School in Stochastic Modeling and Applications in Beijing, the 2011 HKU-HKUST-Stanford Conference in Quantitative Finance in Hong Kong, and the Seventh World Congress of the Bachelier Finance Society in Sydney. The first author acknowledges financial support from a start-up fund at Columbia University, and the second author acknowledges financial support from a start-up fund at the University of Oxford, research funds from the Nomura Centre for Mathematical Finance and from the Oxford-Man Institute of Quantitative Finance, and GRF Grant #CUHK419511.

[†]Department of Industrial Engineering and Operations Research, Columbia University, S. W. Mudd Building, 500 W. 120th Street, New York, NY 10027, US. Email: <xh2140@columbia.edu>.

[‡]Mathematical Institute and Nomura Centre for Mathematical Finance, and Oxford-Man Institute of Quantitative Finance, the University of Oxford, 24-29 St Giles, Oxford OX1 3LB, UK, and Department of Systems Engineering and Engineering Management, the Chinese University of Hong Kong, Shatin, Hong Kong. Email: <zhouxy@maths.ox.ac.uk>.

1 Introduction

The failure of the classical expected utility theory (EUT) to describe many observed human behaviors has motivated economists to develop alternative models of choice. In the past few decades, a substantial amount of research has been conducted on this subject in two directions. In the first direction, economists have attempted to change or relax some of the axioms in EUT, hoping to find satisfactory models of choice. Examples of this approach include Yaari’s *dual theory of choice* (Yaari, 1987) and Quiggin’s *rank-dependent utility* (Quiggin, 1982; Schmeidler, 1989). In the second direction, which is inspired by behavioral psychology, researchers seek to model the processes that lead to choice. The most celebrated accomplishment along this line is Kahneman and Tversky’s *cumulative prospect theory* (Kahneman and Tversky, 1979; Tversky and Kahneman, 1992). Another notable project is Lopes’ security, potential and aspiration theory (or the *SP/A theory*; Lopes, 1987).

Inspired by these models, we propose and formulate in this paper a new model of choice that considers three emotions relevant to decision making: hope, fear and aspirations—while examining a continuous-time financial portfolio choice problem in which an agent, whose preferences are represented by this model of choice, chooses the portfolio that will optimize his payoff at a given terminal time. Although seemingly contradictory psychological states, hope and fear are typically present simultaneously in the same individual. The former is an optimistic anticipation of good situations and the latter is a pessimistic foreboding of bad situations. Aspirations, by contrast, are responsive to the exigencies and opportunities of each decision nexus. We therefore employ a reversed S-shaped probability weighting (or distortion) function to model hope and fear, and a probabilistic constraint to model aspirations. Notably, our model derives fundamental components from both rank-dependent utility theory and SP/A theory. Its connection with these two theories and the motivations behind our model are discussed in detail in Section 2.

In addition to the introduction of a new portfolio choice model, another contribution of this paper is to solve the portfolio choice problem analytically and explicitly. A probability weighting function is a nonlinear transformation applied to the underlying probability measure when risky choices are evaluated. The presence of such a weighting function ruins both the time-consistency that underlies the dynamic programming principle and the concavity of the preference functional that is essential for any optimization problem. For this reason, the classical approaches that have been developed to address expected utility maximization fail to solve the portfolio choice problem introduced in this paper.

A new approach, known as quantile formulation, has recently been developed to overcome the difficulty associated with probability weighting functions; see e.g., Carlier and Dana (2006), Jin and Zhou (2008), and He and Zhou (2011b). This approach involves changing the decision variable from the random variable of the future payoff to the quantile function of the payoff. This change of variable, in general, recovers the concavity of the underlying preference measure, and thus one can employ either calculus of variations or a pointwise maximization/minimization

approach to solve the optimization problem.

In the context of the model presented in this paper, there is an additional probabilistic constraint representing the aspect of aspirations, which can be reformulated naturally in terms of the quantile function. This constraint introduces considerable technical complications, but we are able to derive the optimal solution explicitly after a careful and involved analysis. The main contribution of this paper, however, is the introduction of the three indices—a fear index, a hope index and a lottery-likeness index—that quantify hope, fear and aspirations, as well as a comparative statics analysis that studies the impacts of these emotions on investing behavior. These indices are put forth naturally in the process of solving the underlying portfolio choice problem. The fear index is a quantity that, when sufficiently large, necessitates portfolio insurance in the optimal portfolio. This index is defined in terms of the curvatures (the second- and first-order derivatives) of the probability weighting function $w(p)$ near $p = 1$; thus, it is related to the exaggeration of small probabilities of extremely bad outcomes or, indeed, to the emotion of fear. On the other hand, the hope index is defined through the first-order derivative of $w(p)$ when p is close to zero, which is relevant to the exaggeration of small probabilities of extremely good outcomes or, in psycho-behavioral terms, to the emotion of hope. We find that a higher level of hope makes the agent include more scenarios under the notion of good states. In addition, in choosing optimal portfolios, an agent with a higher level of hope sets his payoff greater in each of the sufficiently good states than another agent with a lower level of hope. Finally, a sufficiently high level of aspirations forces the agent to construct a lottery-type payoff that has a discontinuity with respect to market conditions. More precisely, under good market conditions, characterized as the realization of a set of good states, the agent’s payoff is exceedingly high, but when market conditions are not good, the agent’s payoff is much lower. Because the payoff resulting from a high level of aspirations thus resembles that of a lottery ticket, we introduce a “lottery-likeness index” to capture its essential characteristic.

Decision making behavior can be characterized as risk-averse or risk-seeking. In the model presented in this paper, both fear and a concave utility function lead to risk-averse behavior. Fear, as pointed out earlier, is an overweight on the left tail of a payoff distribution. The concave utility function, by contrast, captures the aversion of the agent to a mean-preserving spread. In the portfolio choice model presented here, these two elements play qualitatively distinct roles in deciding investing behavior. The presence of a certain level of fear determines whether one needs portfolio insurance. The level of insurance needed, however, is unaffected by the level of fear; rather, it is dependent on, among other factors, the utility function that one employs. On the other hand, both hope and aspirations lead to risk-seeking behavior in decision making, yet these elements have qualitatively different impacts on investing behavior in the portfolio selection model. A higher level of hope causes the agent to include more scenarios among “good” states of the world, implying that he will need to take more leverage in order to reap the payoffs from these states. However, hope and aspirations are qualitatively different, and the agent constructs

a lottery-type payoff only when he has an exceedingly high level of aspirations, thus causing him to gamble outright.

After completing a draft of this paper, we came across [Carlier and Dana \(2011\)](#), in which the authors examine a rank-dependent utility maximization problem. Although this work and the present paper share certain features, there are key differences in motivation, scope and implications. First, [Carlier and Dana](#) set out to produce a normative description of a choice model, so their paper considers the rank-dependent utility maximization without the probabilistic constraint that is used to model aspirations in the present paper. By contrast, our goal is to investigate and capture the decision-making role of the three emotions—hope, fear and aspirations—by solving a portfolio choice problem. Second, [Carlier and Dana](#) focus on obtaining some of the necessary and sufficient conditions for an optimal solution, whereas we are able to derive solutions explicitly. This, in turn, facilitates further analysis, such as comparative statics and numerical implementation. Finally, we introduce a quantitative index for each of the emotions considered in the paper, and we discuss the impact that each can have on trading behavior.

The paper is organized as follows. In [Section 2](#) we review the SP/A theory and the rank-dependent utility theory, from which we derive key elements of our model. In [Section 3](#) we pose our portfolio choice problem in continuous time and present its quantile formulation. [Section 4](#) is devoted to studying the feasibility and well-posedness of the problem. The study of feasibility examines whether the investor’s aspirations are too high, relative to his initial capital, for him to achieve. The study of well-posedness, on the other hand, investigates whether the investor would take infinite leverage on risky assets. In [Section 5](#) we solve the portfolio choice problem thoroughly. Along the way, we introduce our indices for hope, fear and aspirations and study the effects of these emotions on investing behavior. In [Section 6](#) we provide an example in which we specify a particular probability weighting function and a power utility function and employ historical U.S. equity and bond data to obtain some numerical results on the optimal solution to the portfolio choice problem. These numerical results confirm the theoretical results obtained in [Section 5](#). Moreover, using this example we compare our model with the EUT portfolio selection model. Finally, [Section 7](#) concludes the paper and the proof of the main theorem is provided in the Appendix.

2 SP/A Theory and Rank-Dependent Utility Theory

In this section, we explain the motivation for our choice of a reversed S-shaped probability weighting function to model hope and fear and of a probabilistic constraint to model aspirations. Let us start from the SP/A theory, which inspired our model.

SP/A theory is a two-factor theory developed by [Lopes \(1987\)](#). It uses both a *dispositional factor* and a *situational factor* to explain risky preferences and choices. The dispositional factor describes the natural motives that dictate individuals’ risk attitudes. In this regard, risk-aversion

appears to be motivated by a desire for *security*, whereas risk-loving is rationalized by a desire to achieve or maximize *potential*. Individuals are found to seek both security and potential when facing risky choices. The situational factor, by contrast, reflects specific needs or opportunities that the individual faces when making choices.

In the decision model proposed by Lopes (1987), the dispositional factor enters into the individual's objective function. A nonnegative prospect (random payoff) X is evaluated as

$$V(X) := \int_0^{+\infty} x d[-w(1 - F_X(x))] \quad (1)$$

where $F_X(\cdot)$ is the cumulative distribution function (CDF) of X . The nonlinear transformation $w(\cdot)$, in Lopes' terms, is called the *decumulative weighting function*, and the integral is in the Lebesgue-Stieltjes sense.¹ In the SP/A theory, $w(\cdot)$ takes the following form:

$$w(z) := \nu z^{q_s+1} + (1 - \nu)[1 - (1 - z)^{q_p+1}], \quad (2)$$

where $0 \leq \nu \leq 1$ and $q_s, q_p \geq 0$. Clearly, z^{q_s+1} and $1 - (1 - z)^{q_p+1}$ are convex and concave functions, respectively; therefore, according to Yaari's theory, they imply risk-aversion and risk-loving, respectively.² The decumulative weighting function w is a mixture of a convex function and a concave function; thus, it seems to represent the individual's (somewhat conflicting) desire for both security and potential. In Lopes' words, individuals stand “*between hope and fear*.” It is worth mentioning that (1) is a natural extension of the definition in Lopes and Oden (1999) that applies to purely discrete prospects.

The situational factor, on the other hand, is modeled by the probabilistic constraint

$$P(X \geq A) \geq \alpha \quad (3)$$

where A is the *aspiration level* and $0 \leq \alpha \leq 1$ is the *confidence level*. See for instance Lopes and Oden (1999). The pair (A, α) represents individuals' aspirations responding to specific circumstances and opportunities. For instance, if an individual has a loan of amount A which will be due soon, he may be forced to set the loan amount to be the aspiration level. On the other hand, if an individual is prohibited from taking excess risk, then he may set up a low aspiration level with a high confidence level to achieve it.³

While it is quite reasonable to model aspirations (the situational factor) via a probabilistic constraint, it is not convincing to us that the dispositional factor should be modeled as (1) with the weighting function $w(\cdot)$ specified as (2). It is, in our view, neither justifiable nor adequate

¹If we define $\tilde{w}(z) := 1 - w(1 - z)$, the *dual* of w , then (1) can be written as $\int_0^{+\infty} x d\tilde{w}(F_X(x))$ which has been used by some authors. Here, we follow the convention in Tversky and Kahneman (1992) to use w instead of $\tilde{w}$.

²See Theorem 2, Yaari (1987).

³Indeed, (3) can also be interpreted as a type of VaR constraint. Such a constraint is popular in the practice of risk management.

to explain the complex psychological interlacement of hope and fear as a mere simple convex combination of two completely opposite attitudes toward risk. Rather, it is more appropriate to attempt to capture *simultaneously* hope as the optimism over extremely satisfactory situations and fear as the pessimism over very poor situations. On the other hand, the use of a linear utility in (1) is also debatable. Linear utility functions are not favored in the risky choice literature. Moreover, we will show later that the preference functional (1) leads to an ill-posed portfolio choice model in the continuous-time setting, mainly due to the linearity of the utility function.

Thus, it is desirable to seek a more suitable preference representation to describe the dispositional factor. We turn therefore to the *rank-dependent utility* (RDU). RDU was introduced by Quiggin (1982) and further developed by Schmeidler (1989). It can explain many paradoxes that EUT has failed to capture, and, at the same time, it provides mathematical tractability. As stated in Starmer (2000), “the rank-dependent model is likely to become more widely used,” because it captures many robust empirical phenomena “in a model which is quite amenable to application within the framework of conventional economic analysis.”

In RDU, a prospect X is evaluated as

$$V(X) := \int_0^{+\infty} u(x) d[-w(1 - F_X(x))] \quad (4)$$

where $w(\cdot)$ is called a *probability weighting/distortion function*. Mathematically, the RDU preference measure (4) generalizes the SP/A counterpart (1) by involving a nonlinear utility function $u(\cdot)$. However, the two measures have distinct economic interpretations. The SP/A preference measure (1) can be regarded as a mixture of two preference measures in Yaari’s dual theory. As we argued earlier, hope and fear are not suitably modeled by this mixture. On the other hand, a rank-dependent utility preference measure, as we will see later, can provide a meaningful characterization of hope and fear.

Abundant work has been done to elicit the utility function and weighting function from experimental data, both using a parameter-based method (Tversky and Fox, 1995; Tversky and Kahneman, 1992) and using a parameter-free method (Abdellaoui, 2000). The typical resulting weighting function is *reversed S-shaped*, showing that small probabilities of both very good and very bad events are overweighted. Meanwhile, the typical utility function is concave, reflecting the fact that people are less favorable disposed toward a risky gamble than to its mean payoff when only intermediate probabilities are involved.

Although RDU is established in a normative way, i.e., via axiomatization, it shows deep psychological intuition. Because the typical weighting function is reversed S-shaped, its high end exhibits convexity, indicating that the individual pays too much attention to the worst outcomes⁴ and hence shows pessimism or *fear*. At the same time, the weighting function is concave at the

⁴Assuming $w(\cdot)$ is differentiable, it follows from (4) that $V(X) = \int_0^{+\infty} w'(1 - F_X(x))u(x)dF_X(x)$, assuming that $F_X(\cdot)$ is continuous. Hence, $w'(\cdot)$ serves as a weight on the utility of the outcome. Clearly the high end of the weighting function corresponds to the low end of the outcome.

low end, reflecting the fact that the individual also pays too much attention to the best outcomes and shows optimism or *hope*. Therefore, RDU—in particular, the reversed S-shaped weighting function—does indeed capture both fear and hope simultaneously.

The following three families of parameterized weighting functions are popular in the literature: the [Tversky and Kahneman \(1992\)](#) weighting function

$$w(p) = \frac{p^\gamma}{(p^\gamma + (1-p)^\gamma)^{1/\gamma}} \quad (5)$$

with $0 < \gamma < 1$; the [Tversky and Fox \(1995\)](#) weighting function

$$w(p) = \frac{\delta p^\gamma}{\delta p^\gamma + (1-p)^\gamma} \quad (6)$$

with $\delta > 0, 0 < \gamma < 1$; and the [Prelec \(1998\)](#) weighting function

$$w(p) = e^{-\delta(-\ln p)^\gamma} \quad (7)$$

with $\delta > 0, 0 < \gamma < 1$. All of these weighting functions are reversed S-shaped.⁵ Estimates of the parameters are available in many papers, such as [Abdellaoui \(2000\)](#), [Wu and Gonzalez \(1996\)](#), [Tversky and Kahneman \(1992\)](#), [Camerer and Ho \(1994\)](#), [Bleichrodt and Pinto \(2000\)](#), and [Abdellaoui, Bleichrodt, and Paraschiv \(2007\)](#).

Inspired by [Jin and Zhou \(2008\)](#) and [Wang \(2000\)](#), we will consider in this paper another class of weighting functions. Parameterized by $(a, b, \bar{z})$, this class of weighting functions is defined as follows:

$$w(z) = \begin{cases} ke^{(a+b)\Phi^{-1}(\bar{z}) + \frac{a^2}{2}} \Phi(\Phi^{-1}(z) + a), & z \leq \bar{z}, \\ A + ke^{\frac{b^2}{2}} \Phi(\Phi^{-1}(z) - b), & z \geq \bar{z}, \end{cases} \quad (8)$$

where $\Phi(\cdot)$ is the CDF of a standard normal random variable, $a, b \geq 0$, and k and A are given as

$$k = \frac{1}{e^{\frac{b^2}{2}} \Phi(-\Phi^{-1}(\bar{z}) + b) + e^{(a+b)\Phi^{-1}(\bar{z}) + \frac{a^2}{2}} \Phi(\Phi^{-1}(\bar{z}) + a)} > 0, \quad A = 1 - ke^{\frac{b^2}{2}},$$

respectively. Essentially, the weighting function (8) is obtained by pasting together smoothly the two one-parameter weighting functions in [Wang \(2000\)](#).⁶ On the other hand, in contrast to [Jin and Zhou \(2008\)](#), where the values of a and b are restricted in a particular range in order to fulfill

⁵The values of the parameters in (5) and (6) must be restricted to certain ranges to make the resulting weighting functions increasing. For instance, [Ingersoll \(2008\)](#) shows that for the Tversky–Kahneman weighting function (5) to be increasing, γ must be larger than 0.28. Luckily, all the estimates of those parameters in the literature lead to increasing weighting functions.

⁶[Wang \(2000\)](#) considers the probability weighting function $w(z) := \Phi(\Phi^{-1}(z) + \alpha)$, parameterized by α . It is straightforward to verify that this function is concave if $\alpha \geq 0$ and convex if $\alpha \leq 0$.

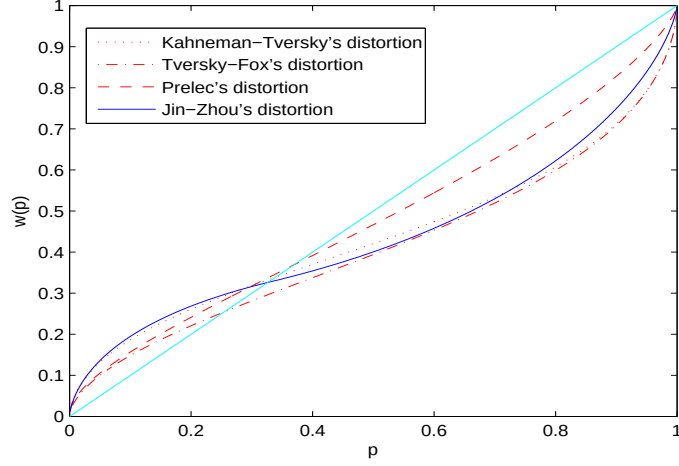


Figure 1: Graph of Kahneman-Tversky weighting function (5), Tversky-Fox weighting function (6), Prelec weighting function (7), Jin-Zhou weighting function (8), and the identity function.

a monotonicity condition therein, we allow a and b to take any nonnegative values because that monotonicity condition is not needed in the present paper.

It is straightforward to compute that

$$w'(z) = \begin{cases} ke^{(a+b)\Phi^{-1}(\bar{z})-a\Phi^{-1}(z)}, & z \leq \bar{z}, \\ ke^{b\Phi^{-1}(z)}, & z \geq \bar{z}. \end{cases}$$

Therefore, $w'(\cdot)$ is decreasing on $(0, \bar{z})$ and increasing on $(\bar{z}, 1)$, and consequently $w(\cdot)$ is reversed S-shaped. In addition, $\bar{z}$ is the inflection point of $w(\cdot)$.

Figure 1 depicts the various weighting functions (5)-(8). We use the parameter values estimated in Tversky and Kahneman (1992), Abdellaoui (2000), and Wu and Gonzalez (1996) for the first three weighting functions, i.e., $\gamma = 0.69$ for (5), $\delta = 0.65$, $\gamma = 0.6$ for (6) and $\gamma = 0.74$ for (7). For the Jin-Zhou weighting function (8), we choose the parameter values so that the resulting weighting function is graphically close to the other three weighting functions.⁷ The specific parameter values are provided in Section 6.

It should be noted that the decumulative weighting function (2) used in the SP/A theory is not, in general, reversed S-shaped. It can generate a reversed S-shaped function with certain parameter values, e.g., $\nu = 0.3$, $q_s = 2$ and $q_p = 6$. However, there is an essential difference between (2) and the classical weighting functions such as (5)-(8). Extremely small probabilities, say 10^{-5} and 10^{-6} , are often indistinguishable to most people; hence, it is reasonable for a weighting function to have infinite sensitivity at both zero and one, like the ones in (5)-(8). However, for the decumulative

⁷Unlike the other three weighting functions (5)-(7), the Jin-Zhou weighting function (8) has not been calibrated to real data. The main reason we use this weighting function here is because it allows for separate investigation of the effects of hope and fear on asset allocation; see Section 6.

weighting function $w(\cdot)$ in (2), $w'(0) = (1 - \nu)(q_p + 1) < \infty$ and $w'(1) = \nu(q_s + 1) < \infty$.

In the rest of this paper, we will formulate and solve our hope, fear and aspiration (HF/A) model, a new portfolio selection model in continuous time. In this model, the objective function is of the form (4), which is taken from RDU with a reversed S-shaped probability weighting function used to capture both hope and fear, whereas a probabilistic constraint of type (3) taken from SP/A theory is used to represent aspirations.⁸

3 Formulation of HF/A Portfolio Choice Model

Let $T > 0$ be the given future time and $(\Omega, \mathcal{F}, (\mathcal{F}_t)_{0 \leq t \leq T}, P)$ be a filtered probability space on which is defined a standard $\mathcal{F}_t$ -adapted n -dimensional Brownian motion $W(t) \equiv (W^1(t), \dots, W^n(t))^\top$ with $W(0) = 0$. It is assumed that $\mathcal{F}_t = \sigma\{W(s) : 0 \leq s \leq t\}$, augmented by all the P -null sets. Here and henceforth, $A^\top$ denotes the transpose of a matrix A , and $a \vee b := \max(a, b)$ for $a, b \in \mathbb{R}$.

We define a continuous-time financial market following Karatzas and Shreve (1998). In this market there are $m + 1$ assets being traded continuously. One of the assets is a risk-free *bank account* whose price process follows

$$S_0(t) = e^{\int_0^t r(u) du}, \quad 0 \leq t \leq T, \quad (9)$$

where the *interest rate* $r(\cdot)$ is a progressively measurable process with

$$\int_0^T |r(t)| dt < \infty, \quad P - \text{a.s.}$$

The other m assets are risky stocks whose price processes $S_i(t)$, $i = 1, \dots, m$ satisfy the following stochastic differential equation (SDE):

$$dS_i(t) = S_i(t) \left[\mu_i(t) dt + \sum_{j=1}^n \sigma_{ij}(t) dW^j(t) \right], \quad t \in [0, T]; \quad S_i(0) = s_i > 0, \quad (10)$$

where $\mu_i(\cdot)$ and $\sigma_{ij}(\cdot)$, named as appreciation rate and volatility rate, respectively, are $\mathcal{F}_t$ -

⁸One may wonder why we do not apply the same probability weighting on the aspiration constraint. To answer this question, let us reiterate the different roles of hope, fear and aspirations in affecting decision making. Hope and fear are two general psychological states that affect the individual's preference over prospects. The probability weighting function is merely a model of the two states. The individual's belief of the prospect X is still represented by the CDF $F_X(\cdot)$, rather than by $1 - w(1 - F_X(\cdot))$. On the other hand, aspirations are responsive to exigencies of each decision nexus and are unrelated to the inherent emotions of hope and fear. Thus, in the probabilistic constraint (3) that models aspirations, we should not apply any probability weighting function on $P(X \geq A)$. That said, applying the probability weighting in (3) would not mathematically change the model since the new constraint would be equivalent to (3) with a revised value of α .

progressively measurable stochastic processes with

$$\int_0^T \left[\sum_{i=1}^m |\mu_i(t)| + \sum_{i=1}^m \sum_{j=1}^n |\sigma_{ij}(t)|^2 \right] dt < +\infty, \quad P - \text{a.s.}$$

Set the excess rate of return process

$$B(t) := (\mu_1(t) - r(t), \dots, \mu_m(t) - r(t))^\top,$$

and define the volatility matrix process $\sigma(t) := (\sigma_{ij}(t))_{m \times n}$. The following basic assumption is imposed on the market parameters throughout this paper:

Assumption 1 *There exists an $\mathcal{F}_t$ -progressively measurable, $\mathbb{R}^n$ -valued process $\theta_0(\cdot)$, the so-called market price of risk, with $Ee^{\frac{1}{2} \int_0^T |\theta_0(t)|^2 dt} < +\infty$ such that*

$$\sigma(t)\theta_0(t) = B(t), \quad P - \text{a.s.}, \quad \text{a.e. } t \in [0, T].$$

Assumption 1 is only slightly stronger than the standard no-arbitrage assumption, and the gap between the two assumptions is due to the additional Novikov condition; see Karatzas and Shreve (1998) for details.

Consider an agent, with an initial endowment $x > 0$ and an investment horizon $[0, T]$, whose total wealth at time $t \geq 0$ is denoted by $X(t)$. Assume that the trading of shares takes place continuously in a self-financing fashion and that there are no transaction costs. Then, the wealth process satisfies (see e.g., Karatzas and Shreve, 1998)

$$dX(t) = r(t)X(t)dt + \pi(t)^\top [B(t)dt + \sigma(t)dW(t)], \quad t \in [0, T]; \quad X(0) = x,$$

where $\pi_i(t)$, $i = 1, 2, \dots, m$, denotes the total market value of the agent's wealth in the i -th asset at time t . The process $\pi(\cdot) \equiv (\pi_1(\cdot), \dots, \pi_m(\cdot))^\top$ is called a *portfolio* if it is $\mathcal{F}_t$ -progressively measurable with

$$\int_0^T |\sigma(t)^\top \pi(t)|^2 dt < +\infty, \quad \text{a.s.},$$

and it is tame (i.e., the corresponding discounted wealth process $X(t)/S_0(t)$ is almost surely bounded from below—although the bound may depend on $\pi(\cdot)$). It is standard in the continuous-time portfolio choice literature that a portfolio is required to be tame so as, among other things, to exclude the notorious doubling strategy.

We now formulate the portfolio choice model featuring hope, fear and aspirations in this financial market. At time 0, given an initial wealth x_0 , the agent seeks the optimal trading strategy such that the terminal wealth/payoff satisfies the probabilistic constraint that represents

aspirations and the preference value of the payoff modeled as in RDU that captures hope and fear is maximized. Therefore, the agent faces the following HF/A portfolio choice problem:

$$\begin{aligned} & \text{Max}_{\pi} \quad \int_0^\infty u(x) d[-w(1 - F_{X(T)}(x))] \\ & \text{subject to} \quad dX(t) = r(t)X(t)dt + \pi(t)^\top [B(t)dt + \sigma(t)dW(t)], \quad t \in [0, T]; \quad X(0) = x_0, \\ & \quad P(X(T) \geq A) \geq \alpha, \quad X(t) \geq 0, \quad \forall t \in [0, T], \end{aligned} \quad (11)$$

where $F_{X(T)}(\cdot)$ is the CDF of $X(T)$ viewed at time 0.

When $w(\cdot)$ is the identity function and the aspiration constraint $P(X(T) \geq A)$ is absent, problem (11) becomes a classical portfolio choice problem under EUT, which has been mainly solved by dynamic programming approach in literature. In this approach, a *class* of problems under the same preference measure at different future times and states (t, x) are considered. Time consistency, which stipulates that the optimal strategy planned today must also be optimal in the future, provides a link among these problems. This leads to the dynamic programming principle, and to HJB equations in Markovian settings, from which we can solve all the problems and, in particular, the one at time 0. With probability weighting, however, our problem (11) is inherently time inconsistent (i.e., the class of problems in the future using the same objective function as in (11) are time inconsistent), for which reason dynamic programming fails.⁹

In literature, there are two ways to address the time-inconsistency. One is to consider a so-called *equilibrium* strategies in lieu of optimal ones, in the context of all the agent's incarnations at different times playing games among themselves; see for instance [Ekeland and Lazrak \(2006\)](#), [Björk, Murgoci, and Zhou \(2011\)](#) and the references therein. The other is to consider *pre-committed* optimal strategies, namely, the agent solves the underlying dynamic optimization problem at time 0 for an optimal strategy and commits himself to follow this strategy in the future. Due to the time-inconsistency of our problem, a pre-committed strategy will not in general solve the portfolio choice problem under the same objective function at a future time. However, pre-committed strategies are still important, since they are frequently applied in practice, sometimes with the help of certain commitment devices. For instance, [Barberis \(2012\)](#) finds that the pre-committed strategy of a casino gambler is a stop-loss one (when the model parameters are in reasonable ranges). Many gamblers indeed follow this strategy by using some commitment devices, such as leaving ATM cards at home or bringing a little money to the casino; see [Barberis \(2012\)](#) for a full discussion.

In this paper, we will study the pre-committed strategies, and will apply the martingale approach to derive them. In this approach, we first determine the optimal terminal payoff and then replicate it using some feasible portfolio. For this purpose we assume

Assumption 2 *The market price of risk is unique, i.e., the market is complete.*

⁹ [Ekeland and Pirvu \(2008\)](#) employ dynamic programming to solve an optimal investment-consumption problem with hyperbolic discounting at a fixed time by constructing a class of future problems having *different* objectives as in the original problem. For our problem (11), it seems impossible to do a similar construction.

With the complete market assumption, we could define the pricing kernel

$$\rho := \exp \left\{ - \int_0^T \left[r(s) + \frac{1}{2} |\theta_0(s)|^2 \right] ds - \int_0^T \theta_0(s)^\top dW(s) \right\}, \quad (12)$$

where θ_0 is the unique market price of risk. Then, the HF/A portfolio choice problem can be reformulated as the following optimization problem:

$$\begin{aligned} \text{Max}_X \quad & \int_0^\infty u(x) d[-w(1 - F_X(x))] \\ \text{subject to} \quad & P(X \geq A) \geq \alpha, \\ & E[\rho X] \leq x_0, \quad X \geq 0, \quad X \text{ is } \mathcal{F}_T \text{ measurable}, \end{aligned} \quad (13)$$

where X represents the terminal payoff of certain portfolio.

The objective function in (13) is not concave due to the presence of the weighting function, so the convex dual method that is employed to solve portfolio selection problems under EUT cannot be applied. Here, we employ a new method—the *quantile formulation*—to cope with this difficulty. To use this method, we need to impose the following technical assumption throughout this paper:

Assumption 3 ρ admits no atom and $E\rho < \infty$.

If $(r(\cdot), \mu(\cdot), \sigma(\cdot))$ is deterministic and $\int_0^T |\theta_0(t)|^2 ds > 0$, then ρ is a lognormal random variable, which satisfies Assumption 3.

For a full account of the theory of quantile formulation, see He and Zhou (2011b). In this theory, roughly speaking, we deal with a generic portfolio choice model satisfying two basic assumptions: law-invariance and “the more the better.” Law-invariance means that the preference measure and all of the constraints other than the budget constraint depend only on the probability distribution of the terminal payoff. “The more the better” means that if the agent has more initial budget, he can achieve higher objective value (see Assumption 2.3 in He and Zhou (2011b) for a precise statement). This assumption is very natural for a sensible economic model. With these two assumptions, He and Zhou (2011b) demonstrated that the portfolio choice problem is equivalent to an optimization problem, i.e., the so-called quantile formulation, in which the optimal quantile function of the terminal payoff is to be found. The advantage of this formulation is that concavity is restored and hence traditional optimization techniques become applicable. Furthermore, there is a simple connection between the optimal solutions to the portfolio choice problem and its quantile formulation. If X^* is the optimal terminal payoff to the portfolio choice problem, then its quantile function is optimal to the quantile formulation. On the other hand, once the optimal quantile function $G^*(\cdot)$ is found, the optimal terminal payoff can be recovered by $X^* := G^*(Z_\rho)$ where $Z_\rho := 1 - F_\rho(\rho)$ is a particular uniformly distributed random variable.¹⁰

¹⁰A general result derived from the quantile formulation is that the optimal payoff must be *anti-comonotonic*

One can easily verify that the HF/A portfolio choice model (13) satisfies the aforementioned two basic assumptions, and hence it has the following quantile formulation:¹¹

$$\begin{aligned} \text{Max}_{G(\cdot)} \quad & U(G(\cdot)) = \int_0^1 u(G(z))w'(1-z)dz \\ \text{subject to} \quad & \int_0^1 F_\rho^{-1}(1-z)G(z)dz \leq x_0. \\ & G(\cdot) \in \mathbb{G}, \quad G((1-\alpha)+) \geq A, \quad G(0+) \geq 0, \end{aligned} \quad (14)$$

where $\mathbb{G}$ is the set of all lower-bounded quantile functions, i.e.,

$$\mathbb{G} = \{G(\cdot) : (0, 1) \rightarrow \mathbb{R}, \text{ nondecreasing, left continuous, and } G(0+) > -\infty\}. \quad (15)$$

Assumption 1–3 will be in force in the reminder of this paper so that we could use quantile formulation to solve the HF/A portfolio choice problem.¹² In the remainder of this paper, we also assume throughout that $0 < \alpha < 1$ and $A \geq 0$,¹³ as well as the following technical assumption on the weighting function $w(\cdot)$, which is satisfied by all the functions in (5)–(8):

Assumption 4 $w(\cdot) : [0, 1] \rightarrow [0, 1]$ is continuous and strictly increasing with $w(0) = 0$, $w(1) = 1$. Furthermore, $w(\cdot)$ is continuously differentiable on $(0, 1)$.

4 Feasibility and Well-posedness

An optimization problem is feasible if it admits at least one solution satisfying *all* the constraints. Feasibility is relevant only to the constraints, not the objective. For our model (13) the aspiration constraint $P(X \geq A) \geq \alpha$ gives rise to the feasibility issue that we should deal with first. It turns out that the feasibility issue can be addressed via a goal reaching problem, which has been solved in the literature.

Proposition 1 *Problem (13) is feasible if and only if $x_0 \geq AE \left[\rho \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}} \right]$. Furthermore, if $x_0 = AE \left[\rho \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}} \right]$, then there is only one feasible solution, given as $X = A \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}}$.*

with respect to the pricing kernel ρ . This property heavily relies on Assumption 3. Examples of violation of this property when Assumption 3 does not hold can be found in Ingersoll (2011).

¹¹For a derivation one may follow the same procedure as described in Section 2 of He and Zhou (2011b).

¹²In the case of an incomplete market and/or in the presence of constraints on portfolios, quantile formulation does not work in general. However, when the investment opportunity set— $r(\cdot)$, $b(\cdot)$ and $\sigma(\cdot)$ —is deterministic, the quantile formulation still works even with conic constraints imposed on portfolios. In this case, one needs to replace ρ in the quantile formulation (14) with the *minimal pricing kernel*. See He and Zhou (2011b, Section 4) for details.

¹³If $\alpha = 0$, then the aspiration constraint is not in force. Therefore, the case in which $\alpha = 0$ can be covered by setting $A = 0$. If $\alpha = 1$, then the aspiration constraint becomes a uniformly lower bound on X , in which case we could shift the terminal payoff X by A and turn the portfolio choice model to a model without the aspiration constraint.

Proof Consider the following goal-reaching problem

$$\begin{aligned} & \underset{X}{\text{Max}} && P(X \geq A) \\ & \text{subject to} && E[\rho X] \leq x_0. \\ & && X \geq 0, \quad X \in \mathcal{F}_T. \end{aligned}$$

Problem (13) is feasible if and only if the optimal value of the above goal-reaching problem is larger than α . From He and Zhou (2011b), Theorem 1, we conclude that if $x_0 \geq AE[\rho]$, then the optimal value is 1. When $x_0 < AE[\rho]$, the optimal value is $F_\rho(c^*)$ where c^* is the value such that $AE[\rho \mathbf{1}_{\{\rho \leq c^*\}}] = x_0$. Therefore, problem (13) is feasible if and only if $x_0 \geq AE[\rho]$, or $x_0 < AE[\rho]$ and $F_\rho(c^*) \geq \alpha$, $AE[\rho \mathbf{1}_{\{\rho \leq c^*\}}] = x_0$ for some $c^* \in \mathbb{R}_+$. Equivalently, (13) is feasible if and only if $x_0 \geq AE[\rho \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}}]$. In particular, if $x_0 = AE[\rho \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}}]$, the goal-reaching problem has a unique optimal solution $A \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}}$ and the optimal value is α . As a result, the only feasible solution to (13) is $X = A \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}}$. ■

Proposition 1 shows that, with exogenously given aspiration level and confidence level, the agent must be sufficiently endowed in order for the level of aspirations to be at least feasible. Put differently, the aspiration level relative to the initial wealth cannot be set too high in the HF/A model. Generally speaking, the triplet (x_0, A, α) must be internally consistent so as to make the model minimally meaningful. In the remainder of this paper, to exclude the infeasibility and the trivial case, we assume that $x_0 > AE[\rho \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}}]$.

The next issue is the well-posedness of the portfolio choice problem. An optimization problem is considered well-posed if its optimal value is finite; otherwise, it is ill-posed. In an ill-posed model, optimality is achieved at the extremal points or boundary of the feasible domain. If the portfolio choice problem is ill-posed, the agent is willing to take as much leverage as possible, leading to excessive risk-taking behavior. For detailed discussions on the ill-posedness issue arising in portfolio choice problems and its economic interpretations, see Jin and Zhou (2008) and He and Zhou (2011a).

We have argued that preference measure (1) proposed in Lopes' SP/A does not capture hope and fear well. Here, we will show that, furthermore, the SP/A theory indeed leads to an ill-posed portfolio choice problem in the continuous-time setting.

Theorem 1 *Let $x_0 > AE[\rho \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}}]$. Assume $u(x) \equiv x$ and $\text{essinf} \rho = 0$. If $\liminf_{z \downarrow 0} w'(z) > 0$, then (13) is ill-posed. In particular, if $w(\cdot)$ is given by (2) with $\nu < 1$, then (13) is ill-posed.*

Proof Let $v(x_0)$ be the optimal value of the problem (13) or (14). Let $\tilde{\mathbb{G}} := \{G(\cdot) \mid G(z) = (A + \tilde{G}(1 - \frac{1-z}{\alpha}))\mathbf{1}_{\{1-\alpha < z < 1\}}, \tilde{G}(\cdot) \in \mathbb{G}, \tilde{G}(0+) \geq 0\} \subset \mathbb{G}$ and consider the following problem:

$$\begin{aligned} \text{Max}_{G(\cdot)} \quad & \int_0^1 G(z)w'(1-z)dz \\ \text{Subject to} \quad & G(\cdot) \in \tilde{\mathbb{G}}, \\ & \int_0^1 F_\rho^{-1}(1-z)G(z)dz \leq x_0. \end{aligned}$$

This problem has a smaller optimal value than (13) because of the smaller feasible set. Rewrite the above problem as

$$\begin{aligned} \text{Max}_{\tilde{G}(\cdot)} \quad & Aw(\alpha) + \alpha \int_0^1 \tilde{G}(z)w'(\alpha(1-z))dz \\ \text{Subject to} \quad & \tilde{G}(\cdot) \in \mathbb{G}, \tilde{G}(0+) \geq 0, \\ & \alpha \int_0^1 F_\rho^{-1}(\alpha(1-z))\tilde{G}(z)dz \leq x_0 - AE \left[\rho \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}} \right]. \end{aligned}$$

This problem is identical to Yaari's dual model as given in He and Zhou (2011b). Recalling He and Zhou (2011b), Theorem 3.4, and the fact that $\liminf_{z \downarrow 0} \frac{w'(\alpha z)}{F_\rho^{-1}(\alpha z)} = \liminf_{z \downarrow 0} \frac{w'(z)}{F_\rho^{-1}(z)} = +\infty$, we conclude that the above problem is ill-posed. $\blacksquare$

Theorem 1 suggests that as long as the agent has hope for extremely satisfactory situations (i.e., $\nu < 1$) while the utility function is linear, he will take as high leverage as possible, leading to an ill-posed problem. The fear of possible catastrophic situations is insufficient to prevent him from taking excessive risky exposures. Therefore, the preference measure (1) in Lopes' SP/A theory is not suitable for portfolio choice problems in the continuous-time setting.¹⁴

This finding further enhances the argument for taking the RDU preference instead of the SP/A one as the model for hope and fear. In the remainder of this paper, we impose the following diminishing marginality on the utility function being used:

Assumption 5 $u(\cdot) : \mathbb{R}_+ \rightarrow \mathbb{R}_+$ is strictly increasing and differentiable. Furthermore, $u'(\cdot)$ is strictly decreasing and satisfies the Inada condition, i.e., $u'(0+) = +\infty$ and $u'(+\infty) = 0$.

5 Solutions in the HF/A Model

In this section we develop the procedure in finding solutions to problem (13) by attacking its quantile formulation (14). Along the way we will introduce, rather naturally, various indices for quantifying the level of hope, fear and aspirations and study their impacts on trading behavior.

¹⁴In a single-period complete market with finite states of the world, the portfolio choice problem under SP/A theory may be well-posed; see Shefrin and Statman (2000).

5.1 Optimal Solution under a Monotonicity Condition

Following the general solution scheme in [He and Zhou \(2011b\)](#), we start with applying the Lagrange dual method to (14). For any Lagrange multiplier $\lambda > 0$, we consider the following problem:

$$\begin{aligned} \text{Max}_{G(\cdot)} \quad & U_\lambda(G(\cdot)) = \int_0^1 [u(G(z))w'(1-z) - \lambda G(z)F_\rho^{-1}(1-z)] dz \\ \text{subject to} \quad & G(\cdot) \in \mathbb{G}, \quad G((1-\alpha)+) \geq A, \quad G(0+) \geq 0. \end{aligned} \quad (16)$$

To simplify the notation, let us denote

$$f(x, z) := u(x)w'(1-z) - \lambda x F_\rho^{-1}(1-z), \quad 0 < z < 1. \quad (17)$$

Then, the objective function in (16) becomes

$$U_\lambda(G(\cdot)) = \int_0^1 f(G(z), z) dz.$$

Define the following function:

$$M(z) := \frac{w'(1-z)}{F_\rho^{-1}(1-z)}, \quad 0 < z < 1, \quad (18)$$

which plays an important role in finding optimal solutions.

Let us first ignore the constraint that $G(\cdot)$ must be a quantile function (its monotonicity being the key requirement), and consider the pointwise maximization of integrand $f(x, z)$ for each fixed z . When $z \in (0, 1-\alpha]$, the maximization problem is $\max_{x \geq 0} f(x, z)$. Clearly, the unique maximizer is $x^* = (u')^{-1}\left(\frac{\lambda}{M(z)}\right)$. When $z \in (1-\alpha, 1)$, the maximization problem is $\max_{x \geq A} f(x, z)$ and the unique maximizer is $x^* = (u')^{-1}\left(\frac{\lambda}{M(z)}\right) \vee A$. This pointwise maximization procedure leads to the introduction of the following function:

$$G_\lambda^*(z) := (u')^{-1}\left(\frac{\lambda}{M(z)}\right) \mathbf{1}_{\{z \leq 1-\alpha\}} + \left[(u')^{-1}\left(\frac{\lambda}{M(z)}\right) \vee A\right] \mathbf{1}_{\{z > 1-\alpha\}}. \quad (19)$$

By this construction, $G_\lambda^*(\cdot)$ automatically satisfies the nonnegativity constraint and the aspiration constraint in (16). If, furthermore, $M(\cdot)$ turns out to be nondecreasing, then $G_\lambda^*(\cdot)$ is nondecreasing and hence a quantile function. In this case, $G_\lambda^*(\cdot)$ is optimal to (16).¹⁵

One can easily check that $M(\cdot)$ is nondecreasing when $w(z) \equiv z$ or in general when $w(\cdot)$ is concave. When $w(z) \equiv z$, there is no probability weighting and the rank-dependent utility degenerates into the classical expected utility. In general, the concavity of $w(\cdot)$ represents a risk-

¹⁵This argument is applied in [Jin and Zhou \(2008\)](#) to solve the gain part problem in a portfolio selection model under the cumulative prospect theory.

seeking attitude (Yaari 1987). In this case, our HF/A model incorporates a risk-averse attitude resulting from the concave utility function and a risk-loving attitude described by the probability weighting function, and an optimal solution is to strike a balance between the two conflicting risk attitudes.¹⁶

We now present the optimal solution to (13) assuming $M(\cdot)$ is nondecreasing. As in standard expected utility maximization problems, we also need the following integrability assumption:

Assumption 6 *There exists $c > 0$ such that for any $\lambda > 0$, $E \left[\rho(u')^{-1} \left(\frac{\lambda \rho}{w'(F_\rho(\rho))} \right) \mathbf{1}_{\{\rho \leq c\}} \right] < +\infty$ and $E \left[u \left((u')^{-1} \left(\frac{\lambda \rho}{w'(F_\rho(\rho))} \right) \right) w'(F_\rho(\rho)) \mathbf{1}_{\{\rho \leq c\}} \right] < +\infty$.*

This assumption can be replaced with a weaker one that the asymptotic elasticity of $u(\cdot)$ is strictly less than one. See Jin, Xu, and Zhou (2007) and Kramkov and Schachermayer (1999) for detailed discussions. In classic expected utility maximization problems, i.e., in the case in which $w(z) \equiv z$, Assumption 6 holds in most of the interesting cases, e.g., when ρ is lognormally distributed and $u(\cdot)$ is a power utility function. In the presence of the weighting function, one could check that this assumption still holds when ρ is lognormally distributed, $u(\cdot)$ is a power utility function and $w(\cdot)$ is taken as (5) and (6), or (7) when $\gamma > \frac{1}{2}$. Having said that, Assumption 6 is not always valid. For instance, if $u(\cdot)$ is a power function, ρ is lognormally distributed and $w(\cdot)$ is taken as (7) with a small γ , it is not difficult to show that Assumption 6 fails.

Theorem 2 *Let Assumptions 5-6 hold, $x_0 > AE[\rho \mathbf{1}_{\{\rho \leq F^{-1}(\alpha)\}}]$, and assume $M(\cdot)$ is nondecreasing. Then, the unique optimal solution to (13) is given as*

$$X^* = \left[(u')^{-1} \left(\frac{\lambda^* \rho}{w'(F_\rho(\rho))} \right) \vee A \right] \mathbf{1}_{\{\rho \leq F^{-1}(\alpha)\}} + (u')^{-1} \left(\frac{\lambda^* \rho}{w'(F_\rho(\rho))} \right) \mathbf{1}_{\{\rho > F^{-1}(\alpha)\}}$$

where λ^* is the value such that the initial budget constraint binds, i.e., $E[\rho X^*] = x_0$.

Proof For each fixed $\lambda > 0$, because $M(\cdot)$ is nondecreasing, $G_\lambda^*(\cdot)$ defined in (19) is also nondecreasing. As a result, $G_\lambda^*(\cdot)$ is optimal to (16). Let

$$\begin{aligned} \mathcal{X}(\lambda) &:= \int_0^1 F_\rho^{-1}(1-z) G_\lambda^*(z) dz \\ &= E \left[\rho \left((u')^{-1} \left(\frac{\lambda \rho}{w'(F_\rho(\rho))} \right) \vee A \right) \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}} \right. \\ &\quad \left. + \rho (u')^{-1} \left(\frac{\lambda \rho}{w'(F_\rho(\rho))} \right) \mathbf{1}_{\{\rho > F_\rho^{-1}(\alpha)\}} \right]. \end{aligned}$$

By Assumption 6 and the fact that $E\rho < \infty$, $\mathcal{X}(\cdot)$ is finite-valued and nonincreasing on $(0, +\infty)$. Furthermore, because ρ has no atom, by the monotone convergence theorem, $\mathcal{X}(\cdot)$ is continuous

¹⁶However, as discussed in Section 2, a concave weighting function does not appropriately capture hope and fear.

on $(0, +\infty)$ and

$$\lim_{\lambda \downarrow 0} \mathcal{X}(\lambda) = +\infty, \quad \lim_{\lambda \uparrow +\infty} \mathcal{X}(\lambda) = AE \left[\rho \mathbf{1}_{\{\rho \leq F^{-1}(\alpha)\}} \right].$$

Therefore, for each $x_0 > AE \left[\rho \mathbf{1}_{\{\rho \leq F^{-1}(\alpha)\}} \right]$, there exists $\lambda^* > 0$ such that $\mathcal{X}(\lambda^*) = x_0$. As discussed in the general solution scheme in [He and Zhou \(2011b\)](#), $G_{\lambda^*}^*(\cdot)$ is optimal to (16) and consequently $X^* := G_{\lambda^*}^*(1 - F_\rho(\rho))$ is optimal to (13). The uniqueness follows easily. $\blacksquare$

5.2 Fear, Portfolio Insurance, and Fear Index

Theorem 2 requires the monotonicity of $M(\cdot)$. However, in this subsection we show that $M(\cdot)$ is *not* nondecreasing for many weighting functions proposed in the literature together with a reasonable pricing kernel ρ . Thus, a different methodology needs to be developed to solve (16) without the monotonicity of $M(\cdot)$.

To start, it is straightforward to check that $M(\cdot)$ is nondecreasing if and only if

$$\frac{w''(z)}{w'(z)} \leq \frac{\bar{F}'(z)}{\bar{F}(z)}, \quad 0 < z < 1 \quad (20)$$

where $\bar{F}(z) := F_\rho^{-1}(z)$, $0 < z < 1$.

The following proposition shows, however, that for the weighting functions in (5)-(8) together with a reasonable pricing kernel, (20) is violated.

Proposition 2 *Suppose ρ is lognormally distributed, i.e.,*

$$F_\rho(x) = \Phi \left(\frac{\ln x - \mu_\rho}{\sigma_\rho} \right)$$

for some μ_ρ and $\sigma_\rho > 0$. For any weighting function in (5)-(7) with $0 < \gamma < 1$, there exists $\varepsilon > 0$ such that

$$\frac{w''(z)}{w'(z)} > \frac{\bar{F}'(z)}{\bar{F}(z)}, \quad 1 - \varepsilon < z < 1. \quad (21)$$

For the weighting function in (8), if $b > \sigma_\rho$, then

$$\frac{w''(z)}{w'(z)} > \frac{\bar{F}'(z)}{\bar{F}(z)}, \quad \bar{z} < z < 1. \quad (22)$$

Proof Because $\bar{F}(z) = F_\rho^{-1}(z) = e^{\mu_\rho + \sigma_\rho \Phi^{-1}(z)}$, we have

$$\frac{\bar{F}'(z)}{\bar{F}(z)} = \frac{\sigma_\rho}{\Phi'(\Phi^{-1}(z))}$$

where $\Phi'(\cdot)$ is the probability density function (PDF) of a standard normal random variable. Thus, for the weighting functions (5)-(7), it is sufficient to show that

$$\frac{w''(\Phi(y))}{w'(\Phi(y))} > \frac{\sigma_\rho}{\Phi'(y)}$$

when y goes to $+\infty$. Noticing $1 - \Phi(y) < \Phi'(y)/y, y > 0$, we can show that

$$\liminf_{y \uparrow +\infty} \frac{w''(\Phi(y))}{w'(\Phi(y))} \Phi'(y) = +\infty$$

for all the weighting functions in (5)-(7), which shows that (21) holds for some $\varepsilon > 0$.

For the weighting function (8), it is straightforward to compute that

$$\frac{w''(z)}{w'(z)} = \begin{cases} -\frac{a}{\Phi'(\Phi^{-1}(z))}, & 0 < z < \bar{z}, \\ \frac{b}{\Phi'(\Phi^{-1}(z))}, & \bar{z} < z < 1. \end{cases}$$

So (22) follows easily. ■

Proposition 2 stipulates that $M(\cdot)$ is not nondecreasing with some common weighting functions and a lognormally distributed pricing kernel. The following theorem shows that in this case the behavior of the optimal solution to (13) is changed drastically when compared with Theorem 2.

Theorem 3 *Let Assumption 5 hold and $x_0 > AE \left[\rho 1_{\{\rho \leq F_\rho^{-1}(\alpha)\}} \right]$. Suppose $w(\cdot)$ is twice differentiable and $\bar{F}(\cdot)$ is differentiable. If there exists $\varepsilon > 0$ such that*

$$\frac{w''(z)}{w'(z)} \geq \frac{\bar{F}'(z)}{\bar{F}(z)}, \quad 1 - \varepsilon < z < 1, \quad (23)$$

then for any optimal solution X^ to (13), it must hold that $\text{essinf } X^* > 0$.*

Proof Let X^* be an optimal solution to (13) and let $G^*(\cdot)$ be its quantile function. Then, from the general theory of quantile formulation, $G^*(\cdot)$ is optimal to (14) and $X^* = G^*(1 - F_\rho(\rho))$ P - a.s.. Therefore, it is sufficient to prove that $G^*(0+) > 0$.

Since $G^*(\cdot)$ is optimal to (14), by convex duality, $G^*(\cdot)$ is also optimal to (16) for some $\lambda^* > 0$. Because $\frac{w''(z)}{w'(z)} \geq \frac{\bar{F}'(z)}{\bar{F}(z)}, 1 - \varepsilon < z < 1$, $M(\cdot)$ is nonincreasing on $(0, \varepsilon)$. Suppose $G^*(0+) = 0$. Let $z_1 := \inf\{z \in (0, 1) \mid G^*(z) > 0\}$. We must have $z_1 < 1$ because $G^*(\cdot) \not\equiv 0$. If $z_1 = 0$, there exists $0 < z_2 < \varepsilon$ such that $G^*(z) < M(z)$ for $z \in (0, z_2]$. Define

$$\tilde{G}(z) := \begin{cases} G^*(z_2) & 0 < z \leq z_2 \\ G^*(z) & z_2 < z < 1. \end{cases}$$

If $z_1 > 0$, we pick $z_3 \in (z_1, 1)$ such that $G^*(z_3) > 0$. Define

$$\tilde{G}(z) := \begin{cases} G^*(z) \vee \delta & 0 < z \leq z_3 \\ G^*(z) & z_3 < z < 1 \end{cases}$$

for some $0 < \delta < \min\{G^*(z_3), \min_{0 < z \leq z_3} M(z)\}$. In both cases, $\tilde{G}(\cdot)$ is feasible to (16). Furthermore, it is straightforward to check that $U_{\lambda^*}(\tilde{G}(\cdot)) > U_{\lambda^*}(G^*(\cdot))$, which is a contradiction. Therefore, we must have $G^*(0+) > 0$ and the proof is complete. $\blacksquare$

This theorem shows that an optimal strategy needs to set a *strictly positive* deterministic floor (essinf $X^* > 0$) in wealth, in sharp contrast to the strategy presented in Theorem 2 when $M(\cdot)$ is monotone.¹⁷ Such a deterministic floor is in line with the *portfolio insurance* commonly practiced in the asset management industry. Portfolio insurance is a risk management strategy by means of which a minimum level of wealth is guaranteed across the investment period. Portfolio insurance has been widely studied in the literature; see, e.g. Basak (1995) and Grossman and Zhou (1996). In most studies, the portfolio insurance level is imposed *exogenously*, and it is not clear how the floor changes with respect to various market parameters. By contrast, our model generates a portfolio insurance level *decisively and endogenously*.¹⁸ Moreover, in the following subsection we will derive an explicit expression for this level.

The key condition that has prompted the need for portfolio insurance is (23), which is expressed in terms of the “curvature” of the probability weighting function near 1. Recall that $w(z)$ when z is close to 1 is relevant to the exaggeration of the (small) probabilities of very bad outcomes; hence, $\frac{w''(z)}{w'(z)}$ when z is near 1 is relevant to fear. Theorem, 3, then implies that portfolio insurance is necessary when the agent is sufficiently fearful, which is quantified as $\frac{w''(z)}{w'(z)}$ exceeding a certain (moving) level when z is near 1.

The above discussion motivates us to define the following *fear index* for a weighting function $w(\cdot)$:

$$\mathcal{I}_w(z) := \frac{w''(z)}{w'(z)}, \quad 0 < z < 1. \quad (24)$$

Note that this index is important and relevant only when z is sufficiently close to 1.

The fear index is clearly an analogue of the Arrow-Pratt measure of absolute risk-aversion for

¹⁷For instance, if $\text{esssup } \rho = +\infty$ and there is no probability weighting, i.e., $w(z) \equiv z$, in which case the agent's preference is dictated by the expected utility theory and $M(\cdot)$ is trivially nondecreasing, then the terminal wealth X^* has no strictly positive floor.

¹⁸The possibility of deriving an endogenous portfolio insurance level from a portfolio choice problem has also been illustrated in several papers. Carlier and Dana (2011) derive the optimal demand for contingent claims for an agent with rank-dependent utility and find that the optimal demand may have a flattening part, suggesting portfolio insurance. The flattening structure has also been discussed in Ingersoll (2011) where the author considers a portfolio choice problem faced by an agent whose preference is modeled by cumulative prospect theory.

a utility function.¹⁹ It can be used to measure the agent’s level of fear. The higher this index, the more convex the weighting function is, and the more fear the agent has. We have shown that this index is critical in deciding the monotonicity of $M(\cdot)$ and the optimal behaviors of the agent following our HF/A model. In particular, by Theorem 3, a sufficiently high level of fear, which is characterized by the fear index exceeding a certain threshold, endogenously necessitates portfolio insurance.

We now compute the fear indices for the weighting functions in (5)-(8). For the Kahneman-Tversky weighting function (5), we have $\mathcal{I}_w(z) \approx 2(1-z)^{-1}$ as $z \uparrow 1$. For the Tversky-Fox weighting function (6), we have $\mathcal{I}_w(z) \approx (1-\gamma)(1-z)^{-1}$ as $z \uparrow 1$. For the Prelec weighting function (7), we have $\mathcal{I}_w(z) \approx (1-\gamma)(-\ln p)^{-1}$ as $z \uparrow 1$. For the Jin-Zhou weighting function (8), we have $\mathcal{I}_w(z) = \frac{b}{\Phi'(\Phi^{-1}(z))}$ as $z \uparrow 1$. Thus, for the Kahneman-Tversky weighting function, the degree of fear is independent of the parameter γ . For the Tversky-Fox and the Prelec weighting functions, a smaller γ leads to a higher degree of fear. For the Jin-Zhou’s weighting function, b measures the degree of fear.

In view of Proposition 2, $M(\cdot)$ is typically *not* nondecreasing. From this point, we impose the following assumption in place of the monotonicity of $M(\cdot)$:

Assumption 7 $M(\cdot)$ is continuously differentiable on $(0, 1)$ and there exists $0 < z_0 < 1$ such that $M(\cdot)$ is strictly decreasing on $(0, z_0)$ and strictly increasing on $(z_0, 1)$. Furthermore, $\lim_{z \uparrow 1} M(z) = +\infty$.

Assumption 7 requires $M(\cdot)$ to be first decreasing and then increasing. From an economics point of view, this requirement aligns with a suitable combination of a reversed S-shaped weighting function and a market variable, namely, the agent is fearful of very bad market conditions and hopeful for very good ones. Not surprisingly, Assumption 7 holds with the weighting functions in (5)-(7) and a lognormally distributed market pricing kernel ρ , as shown in Figures 2-4.²⁰

As for the Jin-Zhou weighting function (8), if ρ is lognormally distributed, i.e., $F_\rho(x) =$

¹⁹Although we have been unable to find an explicit definition of the index (24) for probability weighting in the literature, there are works that imply such an index. For instance, Abdellaoui (2002), Theorem 12, establishes that a weighting function w_2 is more “probabilistic risk averse” than another one w_1 if there exists a continuous, convex and strictly increasing function θ such that $w_2 = \theta \circ w_1$. It is not difficult to show that the latter condition is further equivalent to $\mathcal{I}_{w_2}(z) \geq \mathcal{I}_{w_1}(z)$, $0 < z < 1$. On the other hand, when $u(x) \equiv x$, the RDU (4) can be written as $V(X) = \int_0^\infty w(1 - F_X(x))dx$ by integration by parts. Thus, this preference functional admits the local utility $U(x; F) := -\int_0^x w'(1 - F(y))dy$ according to Machina (1982). The (generalized) Arrow-Pratt index for the local utility, $-U''(x; F)/U'(x; F)$, is $w''(1 - F(x))F'(x)/w'(1 - F(x))$. By Theorem 4 in Machina (1982), the higher the Arrow-Pratt index is, the more risk averse the agent is. Thus, the fear index defined here is also related to the generalized Arrow-Pratt index of a preference functional.

²⁰Here we use the same market data as Mehra and Prescott (1985). More precisely, assume that there is only one stock, e.g., the S&P 500 index. The investment opportunity set (r, b, σ) is estimated from the data of real returns of the U.S. S&P 500 and Treasury Bills during the period 1889-1978: $r = 1\%$, $b = 7\%$ and $\sigma = 16.55\%$. The investment period T is taken as 1 year. The pricing kernel can then be computed from (12). As for the weighting functions, the values of parameters are estimated in Tversky and Kahneman (1992), Abdellaoui (2000), and Wu and Gonzalez (1996). On the other hand, we can in fact prove analytically that Assumption 7 is satisfied for the Prelec weighting function (7).

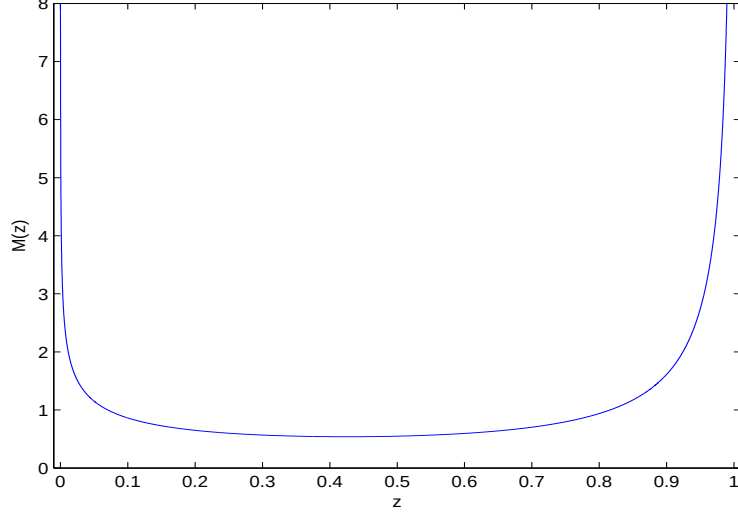


Figure 2: Graph of $M(z)$ with Tversky-Kahneman weighting function.

$\Phi\left(\frac{\ln x - \mu_\rho}{\sigma_\rho}\right)$, then it is easy to compute that

$$M(z) = \begin{cases} ke^{-\mu_\rho + (b - \sigma_\rho)\Phi^{-1}(1-z)}, & 0 < z < 1 - \bar{z}. \\ ke^{(a+b)\Phi^{-1}(\bar{z}) - \mu_\rho - (a + \sigma_\rho)\Phi^{-1}(1-z)}, & 1 - \bar{z} < z < 1. \end{cases}$$

Thus, when $b > \sigma_\rho$, $M(\cdot)$ satisfies Assumption 7 with $z_0 = 1 - \bar{z}$. When $b \leq \sigma_\rho$, $M(\cdot)$ is nondecreasing. In this subsection we are interested in the case in which $b > \sigma_\rho$.

We conclude this subsection by studying the impact of fear on the optimal terminal wealth when the aspiration constraint is absent.

Proposition 3 *Under Assumption 5 and given that $A = 0$, suppose there are two functions, $w_1(\cdot)$ and $w_2(\cdot)$, satisfying Assumptions 6 and 7. Assume*

$$\begin{aligned} w_1(z) &= w_2(z), & 0 < z < 1 - z_0, \\ \mathcal{I}_{w_1}(z) &\geq \mathcal{I}_{w_2}(z) \geq \frac{\bar{F}'(z)}{\bar{F}(z)}, & 1 - z_0 < z < 1. \end{aligned}$$

Let X_1^ and X_2^* be the optimal solutions corresponding to $w_1(\cdot)$ and $w_2(\cdot)$. Then, $X_1^* = X_2^*$.*

Since the proof of Proposition 3 uses a result provided in the following subsection, we defer it to the end of that subsection.

Proposition 3 states that as long as the agent is sufficiently fearful that portfolio insurance becomes necessary, the level of fear affects neither the optimal terminal wealth nor—perhaps more surprisingly—the level of portfolio insurance necessitated! This result can nevertheless be

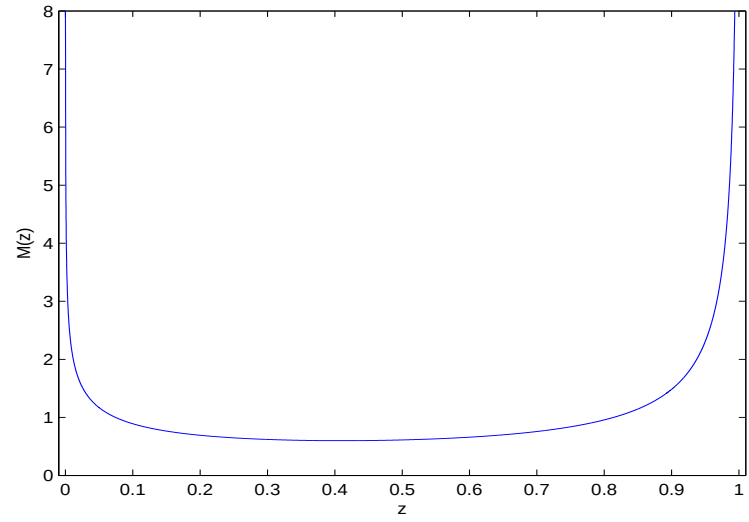


Figure 3: Graph of $M(z)$ with Tversky-Fox weighting function

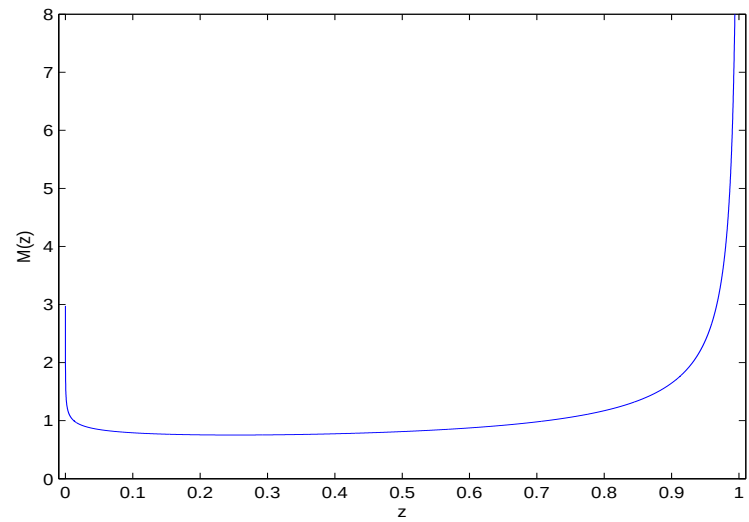


Figure 4: Graph of $M(z)$ with Prelec weighting function.

explained as follows. Recall that fear is described as the tendency of individuals to overweight extremely negative outcomes that occur with small probabilities. In other words, fear is triggered only by the possibility of *catastrophic events*. If the agent is sufficiently fearful, then he will choose to avoid catastrophic events by taking portfolio insurance. Once a portfolio insurance strategy is in place, the degree of fear no longer affects the optimal portfolio because loss due to any catastrophic event will be covered. In this case, the optimal portfolio depends only on the utility function, as well as on the agent's degree of hope and level of aspirations.

5.3 Optimal Solution without the Monotonicity Condition

When $M(\cdot)$ is not nondecreasing, the earlier pointwise maximization argument does not work because $G_\lambda^*(\cdot)$, as defined in (19), is no longer an increasing function. In this case, a new technique is necessary to solve the HF/A portfolio choice problem (13).

Without the monotonicity condition on $M(\cdot)$, the function $G_\lambda^*(\cdot)$ defined by (19) no longer qualifies as a quantile function. Therefore, we have to modify $G_\lambda^*(\cdot)$ in some way so that the resulting function is nondecreasing and, hopefully, a candidate for the optimal solution. To demonstrate the technique that we apply, we first consider the case in which there is no aspiration constraint, i.e., $A = 0$. Then the problem (16) becomes much simpler because the constraint $G((1 - \alpha)+) \geq A$ is removed.²¹

In the following, we denote

$$\bar{G}(z) := (u')^{-1}(\lambda/M(z)), \quad 0 < z < 1. \quad (25)$$

By Assumption 7, $\bar{G}(\cdot)$ is continuous on $(0, 1)$, strictly decreasing on $(0, z_0)$ and strictly increasing on $(z_0, 1)$, and $\lim_{z \uparrow 1} \bar{G}(z) = +\infty$. Recall $f(x, z)$ defined in (17). For any fixed $0 < z < 1$, $f(\cdot, z)$ is strictly increasing on $(0, \bar{G}(z))$ and strictly decreasing on $(\bar{G}(z), +\infty)$.

Define

$$\mathbb{S}_\lambda := \{G(\cdot) \in \mathbb{G} \mid \exists y \in [z_0, 1) \text{ such that } G(z) = \bar{G}(y)\mathbf{1}_{\{0 < z \leq y\}} + \bar{G}(z)\mathbf{1}_{\{y < z < 1\}}, \\ 0 < z < 1\}.$$

We prove in the following proposition that one only needs to consider the type of quantile functions in $\mathbb{S}_\lambda$ in order to solve (16).

Proposition 4 *Let Assumptions 5-7 hold and $A = 0$. For any $G(\cdot) \in \mathbb{G}$ that is feasible to (16), there exists $\tilde{G}(\cdot) \in \mathbb{S}_\lambda$ such that $U_\lambda(\tilde{G}(\cdot)) \geq U_\lambda(G(\cdot))$ and the inequality becomes equality if and only if $G(\cdot) = \tilde{G}(\cdot)$.*

²¹The problem (14) is called a Choquet maximization problem in Jin and Zhou (2008). However, Jin and Zhou (2008) solved the Choquet maximization problem by assuming that $M(\cdot)$ was nondecreasing. So our approach in this paper can also be employed to solve an extension of the Choquet maximization problem and, consequently, that of the portfolio choice problem under CPT in Jin and Zhou (2008).

Proof For any $G(\cdot) \in \mathbb{G}$, let $z_1 := \inf\{z \in (0, z_0] \mid G(z) > \bar{G}(z)\}$ with the convention that $\inf \emptyset = z_0$. Define $z_2 := \inf\{z \in [z_0, 1] \mid \bar{G}(z) > \bar{G}(z_1)\}$. Evidently, $z_0 \leq z_2 < 1$ because $\lim_{z \uparrow 1} \bar{G}(z) = +\infty$. Due to the continuity, $\bar{G}(z_2) = \bar{G}(z_1)$. Define $\tilde{G}(z) := \bar{G}(z_2)\mathbf{1}_{\{0 < z \leq z_2\}} + \bar{G}(z)\mathbf{1}_{\{z_2 < z < 1\}}$. Clearly, $\tilde{G}(\cdot) \in \mathbb{S}_\lambda$. We are going to show that $U_\lambda(\tilde{G}(\cdot)) \geq U_\lambda(G(\cdot))$.

For $0 < z < z_1$, $G(z) \leq G(z_1) \leq \bar{G}(z_1) = \tilde{G}(z) \leq \bar{G}(z)$, where the last inequality is due to the fact that $\bar{G}(\cdot)$ is decreasing on $(0, z_0]$. Therefore, $f(G(z); z) \leq f(\tilde{G}(z); z)$. If $z_1 < z_2$, then $z_1 < z_0$. In the case $z_1 < z < z_2$, $G(z) \geq G(z_1+) \geq \bar{G}(z_1) = \bar{G}(z_2) = \tilde{G}(z) \geq \bar{G}(z)$ and thus $f(G(z), z) \leq f(\tilde{G}(z), z)$. For $z_2 < z < 1$, clearly, $f(G(z), z) \leq f(\bar{G}(z), z) = f(\tilde{G}(z), z)$. Therefore, we have

$$U_\lambda(G(\cdot)) = \int_0^1 f(G(z), z) dz \leq \int_0^1 f(\tilde{G}(z), z) dz = U_\lambda(\tilde{G}(\cdot))$$

and it is easy to see that the inequality becomes equality if and only if $G(\cdot) = \tilde{G}(\cdot)$. ■

In view of Proposition 4, we need only to consider the following problem:

$$\begin{aligned} \text{Max}_{G(\cdot)} \quad & U_\lambda(G(\cdot)) = \int_0^1 [u(G(z))w'(1-z) - \lambda G(z)F_\rho^{-1}(1-z)] dz \\ \text{subject to} \quad & G(\cdot) \in \mathbb{S}_\lambda, \end{aligned} \tag{26}$$

which is essentially a one-dimensional optimization problem, in order to determine the optimal $y \in [z_0, 1)$. To solve this problem, we introduce the following function:

$$\varphi(y) = \int_0^y w'(1-z) dz - M(y) \int_0^y F_\rho^{-1}(1-z) dz, \quad 0 < y < 1. \tag{27}$$

Because $M(\cdot)$ is strictly decreasing on $(0, z_0]$, we have, for any $y \in (0, z_0]$,

$$\frac{\int_0^y w'(1-z) dz}{\int_0^y F_\rho^{-1}(1-z) dz} = \frac{\int_0^y F_\rho^{-1}(1-z) M(z) dz}{\int_0^y F_\rho^{-1}(1-z) dz} > M(y).$$

Thus, we conclude that $\varphi(y) > 0$ on $(0, z_0]$. On the other hand, on $(z_0, 1)$,

$$\begin{aligned} \varphi'(y) &= w'(1-y) - M(y)F_\rho^{-1}(1-y) - M'(y) \int_0^y F_\rho^{-1}(1-z) dz \\ &= -M'(y) \int_0^y F_\rho^{-1}(1-z) dz < 0 \end{aligned}$$

and

$$\lim_{y \uparrow 1} \varphi(y) = -\infty.$$

So there exists a unique root $z^* \in (z_0, 1)$ of $\varphi(\cdot)$ such that $\varphi(z) > 0$ on $(0, z^*)$ and $\varphi(z) < 0$ on $(z^*, 1)$.

Proposition 5 *Let Assumptions 5-7 hold and $A = 0$. Let z^* be the unique root of $\varphi(\cdot)$ defined in (27). Then the unique optimal solution to (16) is*

$$G_\lambda^*(z) = \bar{G}(z^*)\mathbf{1}_{\{0 < z \leq z^*\}} + \bar{G}(z)\mathbf{1}_{\{z^* < z < 1\}}, \quad 0 < z < 1. \quad (28)$$

Proof Let

$$\begin{aligned} V(y) &:= U_\lambda(\bar{G}(y)\mathbf{1}_{\{0 < z \leq y\}} + \bar{G}(z)\mathbf{1}_{\{y < z < 1\}}) \\ &= u(\bar{G}(y)) \int_0^y w'(1-z)dz - \lambda \bar{G}(y) \int_0^y F_\rho^{-1}(1-z)dz \\ &\quad + \int_y^1 u(\bar{G}(z))w'(1-z)dz - \lambda \int_y^1 \bar{G}(z)F_\rho^{-1}(1-z)dz. \end{aligned}$$

By Proposition 4, problem (16) is equivalent to

$$\max_{z_0 \leq y < 1} V(y).$$

We check the first-order derivative of $V(\cdot)$,

$$V'(y) = \frac{\lambda \bar{G}'(y)}{M(y)} \varphi(y),$$

which implies directly that z^* is the unique optimizer. ■

Proposition 5 shows that the optimal solution to (16) is a left truncation of $\bar{G}(\cdot)$. To the left of z^* the optimal quantile function is flat and to the right of z^* it follows $\bar{G}(\cdot)$. Surprisingly, it follows from (27) that the critical point z^* is independent of the multiplier λ —in fact, independent even of the utility function! It depends only on the investment opportunity $F_\rho(\cdot)$ and on the weighting function $w(\cdot)$.

Theorem 4 *Let Assumptions 5-7 hold and $A = 0$. Suppose that $x_0 > 0$. Then (13) has a unique optimal solution*

$$X^* = (u')^{-1} \left(\frac{\lambda^* \rho}{w'(F_\rho(\rho))} \right) \mathbf{1}_{\{\rho \leq c^*\}} + (u')^{-1} \left(\frac{\lambda^* c^*}{w'(F_\rho(c^*))} \right) \mathbf{1}_{\{\rho > c^*\}} \quad (29)$$

where c^* is the unique root of

$$\tilde{\varphi}(x) := x(1 - w(F_\rho(x))) - w'(F_\rho(x)) \int_x^\infty s dF_\rho(s) \quad (30)$$

on $(\text{essinf} \rho, F_\rho^{-1}(1 - z_0))$ and $\lambda^* > 0$ is the value such that $E[\rho X^*] = x_0$.

Proof Let

$$\begin{aligned}\mathcal{X}(\lambda) &:= \int_0^1 F_\rho^{-1}(1-z) G_\lambda^*(z) dz \\ &= (u')^{-1} \left(\frac{\lambda F_\rho^{-1}(1-z^*)}{w'(1-z^*)} \right) \int_0^{z^*} F_\rho^{-1}(1-z) dz \\ &\quad + \int_{z^*}^1 (u')^{-1} \left(\frac{\lambda F_\rho^{-1}(1-z)}{w'(1-z)} \right) F_\rho^{-1}(1-z) dz.\end{aligned}$$

From Assumption 6, it follows that $\mathcal{X}(\lambda)$ is finite and nonincreasing on $(0, +\infty)$. By means of the monotone convergence theorem, recalling that ρ is atomless, one can show that $\mathcal{X}(\cdot)$ is continuous and that

$$\lim_{\lambda \downarrow 0} \mathcal{X}(\lambda) = +\infty, \quad \lim_{\lambda \uparrow +\infty} \mathcal{X}(\lambda) = 0.$$

Therefore, there exists λ^* such that $\mathcal{X}(\lambda^*) = x_0$ and, consequently, $G_{\lambda^*}^*(\cdot)$ is optimal to (14). Then, $X^* := G_{\lambda^*}^*(1 - F_\rho(\rho))$ is optimal to (13). If we let $c^* := F_\rho^{-1}(1 - z^*)$, then X^* is exactly as given in (29). Because z^* is the unique root of $\varphi(\cdot)$ on $(z_0, 1)$, by changing the variable, one can easily check that c^* is the unique root of $\tilde{\varphi}(\cdot)$ on $(\text{essinf } \rho, F_\rho^{-1}(1 - z_0))$. ■

With the optimal terminal wealth profile X^* given by (29), we can divide the states of the world into two classes: the class of “good” states (corresponding to the set $\{\rho \leq c^*\}$) and the class of “bad” states (corresponding to the set $\{\rho > c^*\}$). In bad states of the world, the final wealth has a fixed, deterministic value, $(u')^{-1} \left(\frac{\lambda^* c^*}{w'(F_\rho(c^*))} \right) > 0$, which is irrelevant to which state the world is in. This value is exactly the insured wealth floor. Meanwhile the agent is hopeful for good states, in which the wealth is $(u')^{-1} \left(\frac{\lambda^* \rho}{w'(F_\rho(\rho))} \right)$, the value of which depends on the specific state through ρ .²²

In general, in the presence of the aspiration constraint (i.e. $A > 0$), a similar technique can be applied. We defer the proof to Appendix A and state only the results here.

Introduce the following function:

$$\psi(z) := \frac{\int_0^z w'(1-s) ds}{\int_0^z F_\rho^{-1}(1-s) ds}, \quad 0 < z < 1. \quad (31)$$

On the one hand, compared with $\varphi(\cdot)$ in (27), we conclude that $\psi(z) > M(z)$ on $(0, z^*)$, $\psi(z) < M(z)$ on $(z^*, 1)$ and $\psi(z^*) = M(z^*)$. On the other hand, from $\psi'(z) = -\frac{F_\rho^{-1}(1-z)}{(\int_0^z F_\rho^{-1}(1-s) ds)^2} \varphi(z)$, it follows that $\psi(\cdot)$ is strictly decreasing on $(0, z^*)$ and strictly increasing on $(z^*, 1)$.

²²Because $\lim_{z \uparrow 1} M(z) = +\infty$ due to Assumption 7 and $u'(+\infty) = 0$ due to Assumption 5, we have $\lim_{\rho \rightarrow 0} (u')^{-1} \left(\frac{\lambda^* \rho}{w'(F_\rho(\rho))} \right) = \lim_{\rho \rightarrow 0} (u')^{-1} \left(\frac{\lambda^*}{M(1-F_\rho(\rho))} \right) = +\infty$. Thus, the wealth of the agent in good states is unbounded.

Let

$$x_r := E \left\{ \rho \left[(u')^{-1} \left(\frac{u'(A)M(z^*)\rho}{w'(F_\rho(\rho))} \right) \mathbf{1}_{\{\rho \leq F_\rho^{-1}(1-z^*)\}} + A \mathbf{1}_{\{\rho > F_\rho^{-1}(1-z^*)\}} \right] \right\}, \quad (32)$$

$$x_p := E \left\{ \rho \left[\left((u')^{-1} \left(\frac{u'(A)\psi(1-\alpha)\rho}{w'(F_\rho(\rho))} \right) \vee A \right) \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}} + A \mathbf{1}_{\{\rho > F_\rho^{-1}(\alpha)\}} \right] \right\}. \quad (33)$$

The following theorem provides a complete solution to the HF/A portfolio choice problem (13).

Theorem 5 *Let Assumptions 5 – 7 hold and suppose $x_0 > AE \left[\rho \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}} \right]$. Let c^* be the unique root of $\tilde{\varphi}(\cdot)$ in (30).*

1. Suppose $F_\rho(c^*) < \alpha < 1$.

(a) If $x_0 \geq x_r$, then the unique optimal solution to (13) is

$$X^* := (u')^{-1} \left(\frac{\lambda^* \rho}{w'(F_\rho(\rho))} \right) \mathbf{1}_{\{\rho \leq c^*\}} + (u')^{-1} \left(\frac{\lambda^* c^*}{w'(F_\rho(c^*))} \right) \mathbf{1}_{\{\rho > c^*\}} \quad (34)$$

where $\lambda^* \leq u'(A)M(z^*)$ is the value such that $E[\rho X^*] = x_0$.

(b) If $x_p \leq x_0 \leq x_r$, then the unique optimal solution to (13) is

$$X^* := \left[(u')^{-1} \left(\frac{\lambda^* \rho}{w'(F_\rho(\rho))} \right) \vee A \right] \mathbf{1}_{\{\rho \leq F_\rho^{-1}(1-z_0)\}} + A \mathbf{1}_{\{\rho > F_\rho^{-1}(1-z_0)\}} \quad (35)$$

where $u'(A)M(z^*) \leq \lambda^* \leq u'(A)\psi(1-\alpha)$ is the value such that $E[\rho X^*] = x_0$.

(c) If $x_0 \leq x_p$, then the unique optimal solution to (13) is

$$X^* := \left[(u')^{-1} \left(\frac{\lambda^* \rho}{w'(F_\rho(\rho))} \right) \vee A \right] \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}} + (u')^{-1} \left(\frac{\lambda^*}{\psi(1-\alpha)} \right) \mathbf{1}_{\{\rho > F_\rho^{-1}(\alpha)\}} \quad (36)$$

where $\lambda^* \geq u'(A)\psi(1-\alpha)$ is the value such that $E[\rho X^*] = x_0$.

2. Suppose $0 < \alpha \leq F_\rho(c^*)$, then the unique optimal solution to (13) is

$$X^* := \left[(u')^{-1} \left(\frac{\lambda^* \rho}{w'(F_\rho(\rho))} \right) \vee A \right] \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}} + (u')^{-1} \left(\frac{\lambda^* \rho}{w'(F_\rho(\rho))} \right) \mathbf{1}_{\{F_\rho^{-1}(\alpha) < \rho \leq c^*\}} + (u')^{-1} \left(\frac{\lambda^* c^*}{w'(F_\rho(c^*))} \right) \mathbf{1}_{\{\rho > c^*\}} \quad (37)$$

where λ^* is the multiplier such that $E[\rho X^*] = x_0$.

As in the $A = 0$ case, an optimal strategy divides the states of the world into two or three classes, depending on the parameter values. In the worst states, which are represented by $\{\rho > c^*\}$,

$\{\rho > F_\rho^{-1}(1 - z_0)\}$ or $\{\rho > F_\rho^{-1}(\alpha)\}$, the terminal wealth has a positive floor. In the best states, which are represented by $\{\rho \leq c^*\}$, $\{\rho \leq F_\rho^{-1}(1 - z_0)\}$ or $\{\rho \leq F_\rho^{-1}(\alpha)\}$, the terminal wealth is unbounded from above.

We close this subsection by providing a proof of Proposition 3.

Proof of Proposition 3 It follows from Theorem 5 that the optimal solution to (13) depends, in addition to the pricing kernel ρ , the initial wealth x_0 and the utility function $u(\cdot)$, only on $w(z)$, $0 < z < 1 - z_0$ when $\alpha \leq F_\rho(c^*)$ or when $\alpha > F_\rho(c^*)$ and $x_0 \geq x_r$. Recalling that $A = 0$ implies $\alpha = 0$, we complete our proof. ■

5.4 Hope, Good States of the World, and the Hope Index

We now introduce an index for the agent's level of hope and examine its impact on the optimal payoff. The introduction of the hope index is motivated by the following result:

Proposition 6 *Let Assumption 5 hold. Suppose that there are two functions, $w_1(\cdot)$ and $w_2(\cdot)$, satisfying Assumptions 6 and 7 with the same z_0 . Let c_i^* be the critical point c^* determined as the root of (30) corresponding to $w_i(\cdot)$, $i = 1, 2$, respectively. If*

$$\frac{w'_1(z)}{1 - w_1(z)} \geq \frac{w'_2(z)}{1 - w_2(z)}, \quad 0 < z < 1 - z_0,$$

then $c_1^* \geq c_2^*$.

Proof Let $\varphi_i(\cdot)$ be the function defined in (27) corresponding to $w_i(\cdot)$ and let z_i^* be its unique root in $(z_0, 1)$, $i = 1, 2$. Then, we have $c_i^* = F_\rho^{-1}(1 - z_i^*)$. Now we deduce

$$\begin{aligned} 0 = \varphi_1(z_1^*) &= w'_1(1 - z_1^*) \left[\frac{1 - w_1(1 - z_1^*)}{w'_1(1 - z_1^*)} - \frac{\int_0^{z_1^*} F_\rho^{-1}(1 - z) dz}{F_\rho^{-1}(1 - z_1^*)} \right] \\ &\leq w'_1(1 - z_1^*) \left[\frac{1 - w_2(1 - z_1^*)}{w'_2(1 - z_1^*)} - \frac{\int_0^{z_1^*} F_\rho^{-1}(1 - z) dz}{F_\rho^{-1}(1 - z_1^*)} \right] \\ &= \frac{w'_1(1 - z_1^*)}{w'_2(1 - z_1^*)} \varphi_2(z_1^*). \end{aligned}$$

Thus, we have $\varphi_2(z_1^*) \geq 0$. Because $\varphi_2(\cdot)$ is strictly decreasing in $(z_0, 1)$, we conclude that $z_1^* \leq z_2^*$, and consequently $c_1^* \geq c_2^*$. ■

Inspired by this result, we define

$$\mathcal{H}_w(z) := \frac{w'(z)}{1 - w(z)}, \quad 0 < z < 1 \quad (38)$$

for any given weighting function $w(\cdot)$. Theorem 6 can be restated as the critical point c^* being increasing with respect to $\mathcal{H}_w(z)$ when z is close to 0. According to Theorem 4, c^* divides between good and bad states of the world, and a greater c^* means that the agent includes more scenarios under good states, $\{\rho \leq c^*\}$. Thus, we call $\mathcal{H}_w(z)$ the *hope index*, since this index sheds light on how hope affects investing behavior: the higher the value of the index, the more hopeful the agent is about the future world, and hence the higher leverage he needs to take in order to reap the payoffs from more good states. Note also that the hope index, $\mathcal{H}_w(z)$, is, in terms of the curvature of the weighting function when z is close to 0, consistent with the notion that hope is relevant to the probability weighting of extremely good outcomes.

On the other hand, the following result stipulates that a significantly higher hope index leads to a higher payoff in sufficiently good scenarios than that with a lower hope index.

Proposition 7 *Let Assumption 5 hold and suppose there are two functions, $w_1(\cdot)$ and $w_2(\cdot)$, satisfying Assumptions 6 and 7. Suppose $\text{essinf } \rho = 0$ and the utility function is a power one $u(x) = \frac{x^{1-\eta}-1}{1-\eta}$ where $\eta > 0$. Let the optimal solutions corresponding to $w_1(\cdot)$ and $w_2(\cdot)$ be X_1^* and X_2^* , respectively. If*

$$\lim_{z \downarrow 0} \frac{\mathcal{H}_{w_1}(z)}{\mathcal{H}_{w_2}(z)} = +\infty,$$

then there exists $c > 0$ such that $X_1^ > X_2^*$ on $\{\rho \leq c\}$.*

Proof From Theorem 5, $X_1^* = (u')^{-1} \left(\frac{\lambda_1^* \rho}{w_1'(F_\rho(\rho))} \right)$ and $X_2^* = (u')^{-1} \left(\frac{\lambda_2^* \rho}{w_2'(F_\rho(\rho))} \right)$ for some $\lambda_1^*, \lambda_2^* > 0$ when ρ is sufficiently small. Noticing that

$$\begin{aligned} \lim_{\rho \downarrow 0} \frac{(u')^{-1} \left(\frac{\lambda_1^* \rho}{w_1'(F_\rho(\rho))} \right)}{(u')^{-1} \left(\frac{\lambda_2^* \rho}{w_2'(F_\rho(\rho))} \right)} &= \lim_{\rho \downarrow 0} (u')^{-1} \left(\frac{w_2'(F_\rho(\rho)) \lambda_1^*}{w_1'(F_\rho(\rho)) \lambda_2^*} \right) = (u')^{-1} \left(\lim_{\rho \downarrow 0} \frac{w_2'(F_\rho(\rho)) \lambda_1^*}{w_1'(F_\rho(\rho)) \lambda_2^*} \right) \\ &= (u')^{-1} \left(\lim_{\rho \downarrow 0} \frac{\mathcal{H}_{w_2}(F_\rho(\rho)) \lambda_1^*}{\mathcal{H}_{w_1}(F_\rho(\rho)) \lambda_2^*} \right) = +\infty, \end{aligned}$$

we deduce that there exists $c > 0$ such that $X_1^* > X_2^*$ on $\{\rho \leq c\}$. ■

Let us compute the hope indices for the weighting functions in (5)-(7). For the Tversky-Kahneman weighting (5), we have $\mathcal{H}_w(z) \approx \gamma z^{\gamma-1}$ as $z \downarrow 0$. For the Tversky-Fox weighting function (6), we have $\mathcal{H}_w(z) \approx \delta \gamma z^{\gamma-1}$ as $z \downarrow 0$. For the Prelec weighting function (7), we have $\mathcal{H}_w(z) \approx \delta \gamma \frac{1}{p} (-\ln p)^{\gamma-1} e^{-\delta(-\ln p)^\gamma}$ as $z \downarrow 0$. For each of these three weighting functions, if we have two weighting functions $w_1(\cdot)$ and $w_2(\cdot)$ with $\gamma_1 < \gamma_2$, then $\lim_{z \downarrow 0} \frac{\mathcal{H}_{w_1}(z)}{\mathcal{H}_{w_2}(z)} = +\infty$. In other words, $1 - \gamma$ can measure the hope for good situations. From Proposition 7, the higher the value of $1 - \gamma$ is, the higher the optimal payoff is for good scenarios.

For the Jin-Zhou weighting function (8), the hope index is $\mathcal{H}_w(z) \approx k e^{(a+b)\Phi^{-1}(\bar{z}) - a\Phi^{-1}(z)}$ as $z \downarrow 0$. Furthermore, for any $a_1 > a_2$ and their associated weighting functions, $w_1(\cdot)$ and $w_2(\cdot)$,

we have $\lim_{z \downarrow 0} \frac{\mathcal{H}_{w_1}(z)}{\mathcal{H}_{w_2}(z)} = +\infty$. Therefore, a measures the degree of hope. From Proposition 7, the higher the value of a is, the higher the optimal payoff is for good scenarios.

5.5 Aspirations, Gambles, and the Lottery-Likeness Index

In this subsection we investigate the role that aspirations play in affecting optimal investing behavior, especially when the level of aspirations is exceedingly high. To simplify the discussion, we assume that the utility function is a power function, i.e.,

$$u(x) = \frac{x^{1-\eta} - 1}{1 - \eta} \quad (39)$$

where $\eta > 0$. Nonetheless, all of the qualitative results in this subsection remain true for a general utility function.

Because the utility function is a power function, it is easy to see that the optimal solution to the HF/A portfolio choice problem (13) is proportional to the initial wealth x_0 if the aspiration level is set to be a fixed proportion of the initial wealth. Therefore, in the remainder of this subsection, we assume $x_0 = 1$ without loss of generality, and we consider A to be the aspiration level *relative* to the initial wealth.

Define

$$k_r := \frac{1}{E \left[\left(\frac{\rho}{w'(F_\rho(\rho))} \frac{w'(F_\rho(c^*))}{c^*} \right)^{-\frac{1}{\eta}} \rho \mathbf{1}_{\{\rho \leq c^*\}} + \rho \mathbf{1}_{\{\rho > c^*\}} \right]}, \quad (40)$$

$$k_p := \frac{1}{E \left[\left(\left(\frac{\rho}{w'(F_\rho(\rho))} \psi(1 - \alpha) \right)^{-\frac{1}{\eta}} \vee 1 \right) \rho \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}} + \rho \mathbf{1}_{\{\rho > F_\rho^{-1}(\alpha)\}} \right]}, \quad (41)$$

$$k_u := \frac{1}{E \left[\rho \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}} \right]}. \quad (42)$$

To study the impact of aspirations on trading behavior, we first reproduce Theorem 5 when $u(\cdot)$ is a power function.

Corollary 1 *Suppose $u(\cdot)$ is given in (39), $x_0 = 1$ and let Assumptions 6 and 7 hold. Let c^* be the unique root of $\tilde{\varphi}(\cdot)$ in (30). If $A > k_u$, the problem (13) is infeasible. Otherwise, we have the following assertions:*

1. Suppose $F_\rho(c^*) < \alpha < 1$.

(a) When $0 \leq A \leq k_r$, the unique optimal solution to (13) is

$$X^* = k_r \left(\frac{w'(F_\rho(c^*))}{c^*} \right)^{-\frac{1}{\eta}} \left[\left(\frac{\rho}{w'(F_\rho(\rho))} \right)^{-\frac{1}{\eta}} \mathbf{1}_{\{\rho \leq c^*\}} + \left(\frac{c^*}{w'(F_\rho(c^*))} \right)^{-\frac{1}{\eta}} \mathbf{1}_{\{\rho > c^*\}} \right].$$

(b) When $k_r \leq A \leq k_p$, the unique optimal solution to (13) is

$$X^* = A \left[\left(\left(\lambda_1(A) \left(\frac{\rho}{w'(F_\rho(\rho))} \right)^{-\frac{1}{\eta}} \right) \vee 1 \right) \mathbf{1}_{\{\rho \leq F_\rho^{-1}(1-z_0)\}} + \mathbf{1}_{\{\rho > F_\rho^{-1}(1-z_0)\}} \right]$$

where $\lambda_1(A)$ is the unique number in the interval $\left[\psi(1-\alpha)^{-\frac{1}{\eta}}, \left(\frac{w'(F_\rho(c^*))}{c^*} \right)^{-\frac{1}{\eta}} \right]$ such that

$$E \left[\left(\left(\lambda \left(\frac{\rho}{w'(F_\rho(\rho))} \right)^{-\frac{1}{\eta}} \right) \vee 1 \right) \rho \mathbf{1}_{\{\rho \leq F_\rho^{-1}(1-z_0)\}} + \rho \mathbf{1}_{\{\rho > F_\rho^{-1}(1-z_0)\}} \right] = \frac{1}{A}.$$

(c) When $k_p \leq A \leq k_u$, the unique optimal solution to (13) is

$$X^* = A \left[\left(\left(\lambda_2(A) \left(\frac{\rho}{w'(F_\rho(\rho))} \right)^{-\frac{1}{\eta}} \right) \vee 1 \right) \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}} + \lambda_2(A) \left(\frac{1}{\psi(1-\alpha)} \right)^{-\frac{1}{\eta}} \mathbf{1}_{\{\rho > F_\rho^{-1}(\alpha)\}} \right]$$

where $\lambda_2(A)$ is the unique number in the interval $[0, \psi(1-\alpha)^{-\frac{1}{\eta}}]$ such that

$$E \left[\left(\left(\lambda \left(\frac{\rho}{w'(F_\rho(\rho))} \right)^{-\frac{1}{\eta}} \right) \vee 1 \right) \rho \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}} + \lambda \left(\frac{1}{\psi(1-\alpha)} \right)^{-\frac{1}{\eta}} \rho \mathbf{1}_{\{\rho > F_\rho^{-1}(\alpha)\}} \right] = \frac{1}{A}.$$

2. Suppose $0 < \alpha \leq F_\rho(c^*)$.

(a) When $0 \leq A \leq \left(\frac{w'(F_\rho(c^*))F_\rho^{-1}(\alpha)}{c^*w'(\alpha)} \right)^{-\frac{1}{\eta}} k_r$, the unique optimal solution to (13) is

$$X^* = k_r \left(\frac{w'(F_\rho(c^*))}{c^*} \right)^{-\frac{1}{\eta}} \left[\left(\frac{\rho}{w'(F_\rho(\rho))} \right)^{-\frac{1}{\eta}} \mathbf{1}_{\{\rho \leq c^*\}} + \left(\frac{c^*}{w'(F_\rho(c^*))} \right)^{-\frac{1}{\eta}} \mathbf{1}_{\{\rho > c^*\}} \right].$$

(b) When $\left(\frac{w'(F_\rho(c^*))F_\rho^{-1}(\alpha)}{c^*w'(\alpha)} \right)^{-\frac{1}{\eta}} k_r \leq A \leq k_u$, the unique optimal solution to (13) is

$$\begin{aligned} X^* := & A \left[\left(\left(\lambda_3(A) \left(\frac{\rho}{w'(F_\rho(\rho))} \right)^{-\frac{1}{\eta}} \right) \vee 1 \right) \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}} \right. \\ & \left. + \lambda_3(A) \left(\frac{\rho}{w'(F_\rho(\rho))} \right)^{-\frac{1}{\eta}} \mathbf{1}_{\{F_\rho^{-1}(\alpha) < \rho \leq c^*\}} + \lambda_3(A) \left(\frac{c^*}{w'(F_\rho(c^*))} \right)^{-\frac{1}{\eta}} \mathbf{1}_{\{\rho > c^*\}} \right] \end{aligned}$$

where $\lambda_3(A)$ is the unique number in the interval $\left[0, \left(\frac{w'(\alpha)}{F_\rho^{-1}(\alpha)} \right)^{-\frac{1}{\eta}} \right]$ such that

$$\begin{aligned} E \left\{ \left[\lambda \left(\frac{\rho}{w'(F_\rho(\rho))} \right)^{-\frac{1}{\eta}} \vee A \right] \rho \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}} + \lambda \left(\frac{\rho}{w'(F_\rho(\rho))} \right)^{-\frac{1}{\eta}} \rho \mathbf{1}_{\{F_\rho^{-1}(\alpha) < \rho \leq c^*\}} \right. \\ \left. + \lambda \left(\frac{c^*}{w'(F_\rho(c^*))} \right)^{-\frac{1}{\eta}} \rho \mathbf{1}_{\{\rho > c^*\}} \right\} = \frac{1}{A}. \end{aligned}$$

Proof Corollary 1 is the direct consequence of Theorem 5, taking into consideration the special form of $u(\cdot)$. ■

From Corollary 1-1(a), we can see that k_r is the portfolio insurance level when there are no aspirations. It is clear that the portfolio insurance level increases with respect to the relative risk aversion η .

Aspirations are responsive to immediate, specific needs or opportunities of the agent's decision nexus. Intuitively, a higher level of aspirations will drive the agent to be more aggressive and to take on more risk. An extremely high level of aspirations will *force* the agent to exhibit a lottery-buying type of investment behavior, namely, gambling on a big gain with a risk of losing everything. Since $X^* = G^*(1 - F_\rho(\rho)) \geq G^*(1 - \alpha) \geq A$, i.e., the aspirations are achieved when $\rho \leq F_\rho^{-1}(\alpha)$, we can regard $\{\rho \leq F_\rho^{-1}(\alpha)\}$ as “winning states.” Similarly, $\{\rho > F_\rho^{-1}(\alpha)\}$ consists of “losing states” where the aspirations are not met. Motivated by this, we define the following *lottery-likeness index*:

$$\mathcal{L}(A) = \frac{\text{essinf}(X^* \mid \rho \leq F_\rho^{-1}(\alpha))}{\text{esssup}(X^* \mid \rho > F_\rho^{-1}(\alpha))}, \quad (43)$$

which gives the ratio between the worst winning payoff and the best losing payoff. For a payoff like that of a lottery, the value of this index is expected to be extremely large.²³

Theorem 6 Suppose $u(\cdot)$ is given in (39), $x_0 = 1$ and let Assumptions 6 and 7 hold. Let c^* be the unique root of $\tilde{\varphi}(\cdot)$ in (30). Suppose $A \leq k_u$, and let $\lambda_2(A)$ and $\lambda_3(A)$ be defined in Corollary 1. When $F_\rho(c^*) < \alpha < 1$, we have

$$\mathcal{L}(A) = \begin{cases} 1, & 0 \leq A \leq k_p, \\ \frac{\psi(1-\alpha)^{-\frac{1}{\eta}}}{\lambda_2(A)}, & k_p < A < k_u. \end{cases}$$

When $0 < \alpha \leq F_\rho(c^*)$, we have

$$\mathcal{L}(A) = \begin{cases} 1, & 0 \leq A \leq \left(\frac{F_\rho^{-1}(\alpha)w'(F_\rho(c^*))}{w'(\alpha)c^*} \right)^{-\frac{1}{\eta}} k_r, \\ \frac{1}{\lambda_3(A)} \left(\frac{w'(\alpha)}{F_\rho^{-1}(\alpha)} \right)^{-\frac{1}{\eta}}, & \left(\frac{F_\rho^{-1}(\alpha)w'(F_\rho(c^*))}{w'(\alpha)c^*} \right)^{-\frac{1}{\eta}} k_r \leq A < k_u. \end{cases}$$

Furthermore, $\mathcal{L}(A)$ is increasing in A , and

$$\lim_{A \uparrow k_u} \mathcal{L}(A) = +\infty.$$

Proof The expression of $\mathcal{L}(A)$ can be calculated directly from Corollary 1. Furthermore, it is

²³For example, if a jackpot in a mark six lottery is regarded as the winning state and all the other prizes (including no prize) as the losing states, then the corresponding index value (43) is extremely large.

easy to see that $\lambda_2(A)$ and $\lambda_3(A)$ are decreasing in A and go to $+\infty$ as A goes to k_u . Consequently, $\mathcal{L}(A)$ is increasing in A and $\lim_{A \uparrow k_u} \mathcal{L}(A) = +\infty$. $\blacksquare$

The fact that $\mathcal{L}(A) > 1$ for sufficiently large A indicates that there is a discontinuity of the terminal wealth $X^* \equiv X^*(\rho)$ at $\rho = F_\rho^{-1}(\alpha)$. This discontinuity suggests that the least payoff in “winning states” is greater than the best payoff in “losing states,” which is a characteristic feature of payoffs of lottery tickets. Theorem 6 confirms that a higher level of aspirations will make the trading behavior more like that of buying into a lottery, which in turn justifies the introduction of the lottery-likeness index to quantify the impact of aspirations on trading behavior.

Recall that in our HF/A model, there are two factors that may drive the agent to be risk-taking: a high level of aspirations and a high level of hope, the latter being reflected by a large value of the hope index or by a steep curvature of $w(z)$ when z is close to 0. It is a natural question whether a high level of hope will also induce lottery-buying behavior. Assume for simplicity of the discussion that $A = 0$. According to Theorem 4, the winning states are given by $\{\rho \leq c^*\}$, while the losing states are given by $\{\rho > c^*\}$. Then, an analogous definition of the index (43) is

$$\frac{\text{essinf}(X^* \mid \rho \leq c^*)}{\text{esssup}(X^* \mid \rho > c^*)} \equiv 1.$$

In other words, even an extremely high level of hope will not lead to lottery-buying behavior.

6 Numerical Experiments

In this section, we report the results of our numerical experiments with the aim of demonstrating the analytical findings in the previous section. In our experiments, we assume a constant interest rate r , a constant stock expected return rate μ and a constant stock volatility rate σ . As a result, the market price of risk is $\theta := \sigma^{-1}(\mu - r)$. Given a terminal time T , the pricing kernel ρ is lognormally distributed, i.e., $\ln \rho$ is a normal random variable with the mean and standard deviation

$$\mu_\rho = -\left(r + \frac{\|\theta\|^2}{2}\right)T, \quad \sigma_\rho = \|\theta\|\sqrt{T}.$$

We use a power utility in (39) and the Jin-Zhou weighting function (8). We use the Jin-Zhou weighting instead of the ones in (5)-(7) for the following two reasons: First, as observed in Sections 5.2 and 5.4, for the Tversky-Kahneman weighting function, the fear index is independent of the parameter γ , whereas the hope index depends on γ . For the Tversky-Fox and Prelec weighting functions, both the degrees of hope and fear are measured by the same parameter γ . The Jin-Zhou weighting function is therefore the only one that has *separate* parameters, a and b , to measure the degrees of hope and fear, respectively. This allows us to study the impact of hope and fear on asset allocation separately by varying one parameter and fixing the other. Secondly, by using

the Jin-Zhou weighting function, we can compute the optimal portfolios explicitly. As we will see shortly, such explicit solutions facilitate a comparison between the optimal allocation to risky assets in our model and that in the classical expected utility model.

We use the data set from [Mehra and Prescott \(1985\)](#) to provide estimates for r , μ and σ . Calibrating to the real returns of Treasury Bills for the period 1889-1978, we set $r = 1\%$. We assume only one risky asset, which can be considered to be the market portfolio, and we use the S&P 500 index as a proxy for the market portfolio. Using the real returns of the S&P 500 index for the period 1889-1978, we set $\mu = 7\%$ and $\sigma = 15.34\%$. The terminal time T , which measures how frequently investors evaluate their portfolios, is set at one year in light of the argument presented in [Benartzi and Thaler \(1995\)](#). As a result, we have $\theta = 39.11\%$, $\mu_\rho = -6.65\%$ and $\sigma_\rho = 39.11\%$.

[Lucas \(1994\)](#) claims that the reasonable range of the relative risk aversion η should be between 1 and 2.5. Thus, we set $\eta = 1.5$ by taking a value in the middle of the range. Because the Jin-Zhou weighting function has not been calibrated to real data in the literature, we choose parameter values such that the resulting weighting function is graphically close to the weighting functions (5)-(7), which are based on estimates available in the literature. We set $\bar{z} = \frac{1}{3}$ because it has been reported in the literature that the inflection point of the estimated weighting functions is near $1/3$; see e.g., [Abdellaoui \(2000\)](#). On the other hand, we set $a = 3\sigma_\rho$ and $b = 2.2\sigma_\rho$. The resulting weighting function is graphed in Figure 1 in Section 2 in comparison with the three classical weighting functions. These values of a and b are fixed as a benchmark. Recall that a measures the degree of hope and b measures the degree of fear. When studying the impact of hope and fear on asset allocation, we will vary the values of a and b , respectively.

Finally, we set the initial wealth $x_0 = 1$ without loss of generality.

Figures 5 and 6 below show the optimal terminal payoffs as functions of the pricing kernel with different confidence and aspiration levels. In Figure 5, α is set at 0.5 and the aspiration level A is given as 0.8, 0.91 and 1.5. In Figure 6, α is 0.1 and the aspiration level A is given as 1 and 3.4 respectively. We see that in both figures the functions with high aspiration levels are discontinuous at points that divide winning and losing states, suggesting that the agent is forced to construct a lottery-type payoff. This observation is further confirmed in Figure 7, where we vary the aspiration level A from 0 to k_u and plot the lottery-likeness index $\mathcal{L}$ (again, we consider two confidence levels α , which in this case are set at 0.5 and 0.1, respectively). As predicted by Theorem 6, this index increases when the aspiration level goes up, and it increases sharply when the aspiration level becomes sufficiently high.

Next, we study the impact of hope and fear by varying the values of a and b while fixing the aspiration level at $A = 0$. Figure 8 provides the optimal solutions with a at 0, $2\sigma_\rho$, $3\sigma_\rho$ and $4\sigma_\rho$, respectively, while b is set at $2.2\sigma_\rho$. Figure 9 depicts the optimal payoffs with b at $1.5\sigma_\rho$, $2.2\sigma_\rho$, $3\sigma_\rho$ and $4\sigma_\rho$ while $a = 3\sigma_\rho$. From these two cross-sections, we can see that a higher a or a lower b leads to a lower portfolio insurance level, as well as to a higher payoff in good scenarios.²⁴ To

²⁴Proposition 3 shows that the optimal solution does not depend on the degree of fear if the investor is already

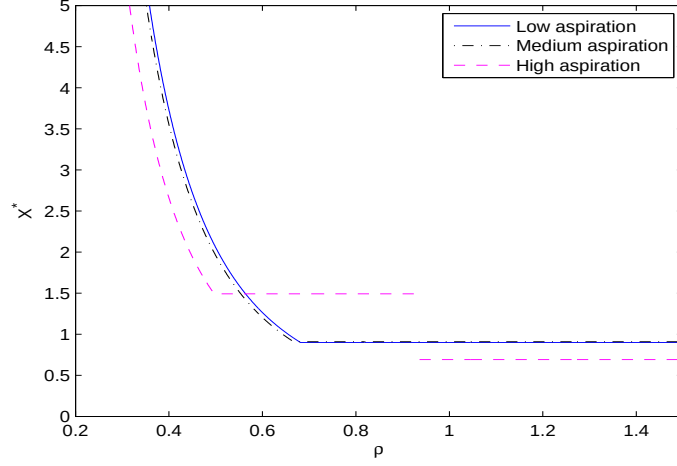


Figure 5: Graph of optimal solution X^* as a function of ρ with different aspiration levels A and a given confidence level $\alpha > F_\rho(c^*)$.

further investigate the impact of hope and fear on portfolio insurance, we take different values of a and b and compute the corresponding portfolio insurance levels. The results are shown in Figure 10. We can see that the portfolio insurance level increases when b increases or when a decreases. Interestingly, the portfolio insurance level is almost always higher than 75 percent of the initial wealth, and, with a high value of b and a low value of a , it is over 95 percent of the initial wealth.

Finally, we set out to find the optimal dynamic portfolio and compare our model with the classic expected utility maximization model. The particular structure of the Jin-Zhou weighting function allows us to calculate the optimal portfolio explicitly. Recall that r is the constant risk-free rate and θ is the constant market price of risk. Let

$$\rho(t) := e^{-(r + \frac{1}{2}\|\theta\|^2)t - \theta^\top W(t)}, \quad 0 \leq t \leq T.$$

Recall that $\Phi(\cdot)$ and $\Phi'(\cdot)$ are the CDF and PDF of the standard normal, respectively. The following theorem provides the optimal portfolio and optimal wealth process under this market setting.²⁵

Theorem 7 Suppose $A = 0$, $u(\cdot)$ is given by (39), $x_0 = 1$ and $w(\cdot)$ is given in (8) with $a \geq 0, b \geq \sigma_\rho$. Let c^* be the unique root of $\tilde{\varphi}(\cdot)$ in (30). Then the optimal wealth process and the

sufficiently fearful. This result only holds true when the probability weighting function does not change at the high end, i.e., when the investor's degree of hope is fixed. Here, when the value of b changes, the whole shape of the probability weighting function changes, leading to a change in the optimal solution.

²⁵In this theorem we assume $A = 0$ for ease of exposition in order to be able to compare our results with those of the utility maximization model. We can also compute the optimal portfolio explicitly with a general aspiration level A .

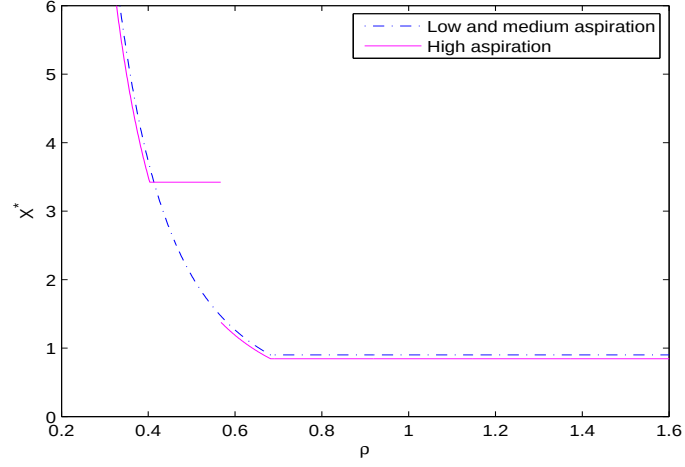


Figure 6: Graph of optimal solution X^* as a function of ρ with different aspiration levels A and a given confidence level $\alpha < F_\rho(c^*)$.

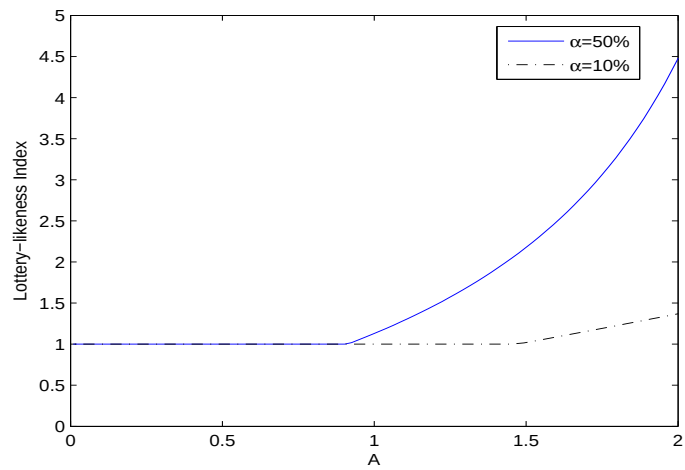


Figure 7: Increase in lottery-likeness index $\mathcal{L}$ with respect to aspiration level A .

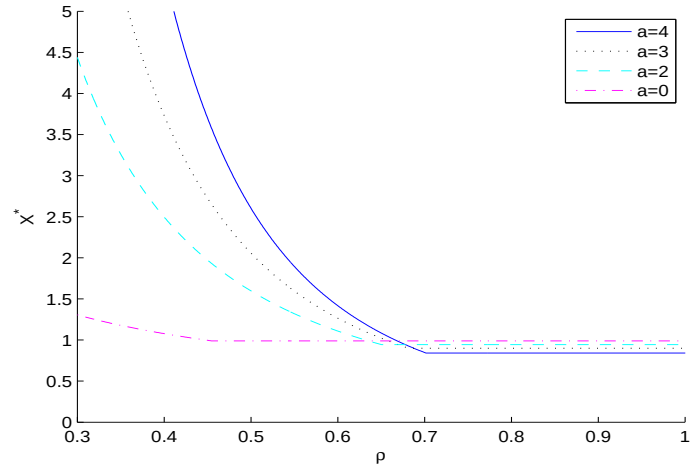


Figure 8: Optimal solution X^* for different values of a without aspiration.

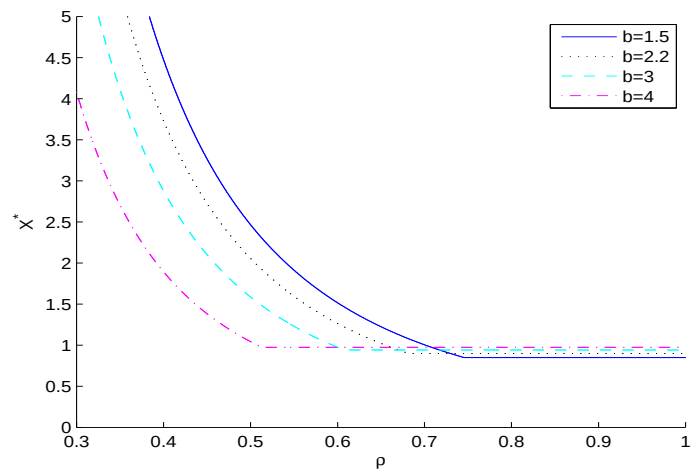


Figure 9: Optimal solution X^* for different values of b without aspiration.

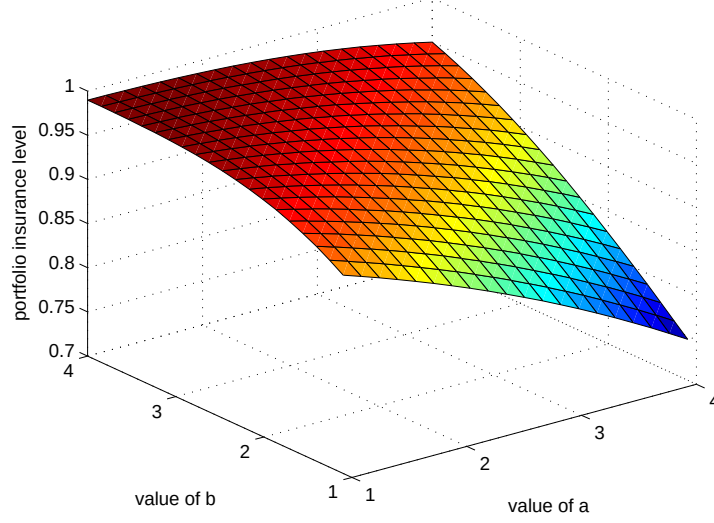


Figure 10: Portfolio insurance level for different values of a and b and zero aspiration level.

corresponding portfolio process are given as

$$X(t) = \Gamma(t, \rho(t)), \quad \pi(t) = \Delta(t, \rho(t)) \frac{1}{\eta} (\sigma^\top)^{-1} \theta X(t), \quad 0 \leq t < T, \quad (44)$$

where

$$\begin{aligned} \Gamma(t, \rho) := & k_r \left[(c^*)^{\frac{1}{\eta} \left(\frac{a}{\sigma_\rho} + 1 \right)} e^{\left(\left(\frac{1}{\eta} \left(\frac{a}{\sigma_\rho} + 1 \right) - 1 \right) r + \frac{1}{\eta} \left(\frac{a}{\sigma_\rho} + 1 \right) \left(\frac{1}{\eta} \left(\frac{a}{\sigma_\rho} + 1 \right) - 1 \right) \frac{\|\theta\|^2}{2} \right) (T-t)} \rho^{-\frac{1}{\eta} \left(\frac{a}{\sigma_\rho} + 1 \right)} \right. \\ & \times \Phi \left(\frac{\ln c^* + \left(r + \left(\frac{1}{\eta} \left(\frac{a}{\sigma_\rho} + 1 \right) - \frac{1}{2} \right) \|\theta\|^2 \right) (T-t) - \ln \rho}{\|\theta\| \sqrt{T-t}} \right) \\ & \left. + e^{-r(T-t)} \left(1 - \Phi \left(\frac{\ln c^* + \left(r - \frac{1}{2} \|\theta\|^2 \right) (T-t) - \ln \rho}{\|\theta\| \sqrt{T-t}} \right) \right) \right], \end{aligned}$$

and

$$\Delta(t, \rho) := \left(\frac{a}{\sigma_\rho} + 1 \right) \left[1 - \frac{k_r e^{-r(T-t)} \left(1 - \Phi \left(\frac{\ln c^* + \left(r - \frac{1}{2} \|\theta\|^2 \right) (T-t) - \ln \rho}{\|\theta\| \sqrt{T-t}} \right) \right)}{\Gamma(t, \rho)} \right].$$

Furthermore, for each fixed $t < T$, $\Gamma(t, \rho)$ and $\Delta(t, \rho)$ are strictly decreasing in ρ and

$$\lim_{\rho \downarrow 0} \Gamma(t, \rho) = +\infty, \quad \lim_{\rho \uparrow \infty} \Gamma(t, \rho) = e^{-r(T-t)} k_r, \quad \lim_{\rho \downarrow 0} \Delta(t, \rho) = \frac{a}{\sigma_\rho} + 1, \quad \lim_{\rho \uparrow \infty} \Delta(t, \rho) = 0. \quad (45)$$

Proof It is easy to compute that

$$\frac{w'(F_\rho(x))}{x} = ke^{(a+b)\Phi^{-1}(1-z_0)+a\frac{\mu\rho}{\sigma_\rho}x - \left(\frac{a}{\sigma_\rho}+1\right)}, \quad x \leq F_\rho^{-1}(\bar{z}).$$

Thus, from Corollary 1-1(a), the optimal terminal payoff is given as

$$X^* = k_r(c^*)^{\frac{1}{\eta}\left(\frac{a}{\sigma_\rho}+1\right)} \rho^{-\frac{1}{\eta}\left(\frac{a}{\sigma_\rho}+1\right)} \mathbf{1}_{\{\rho \leq c^*\}} + k_r \mathbf{1}_{\{\rho > c^*\}}.$$

By Theorem E.1 in Jin and Zhou (2008), the portfolio replicating $\rho^{-\frac{1}{\eta}\left(\frac{a}{\sigma_\rho}+1\right)} \mathbf{1}_{\{\rho \leq c^*\}}$ and its wealth process are

$$\begin{aligned} \pi_1(t) &= \left[\frac{1}{\eta} \left(\frac{a}{\sigma_\rho} + 1 \right) X_1(t) + \frac{1}{\|\theta\|\sqrt{T-t}\rho(t)} (c^*)^{1-\frac{1}{\eta}\left(\frac{a}{\sigma_\rho}+1\right)} \Phi' \left(\frac{\ln c^* + \left(r + \frac{\|\theta\|^2}{2}\right)(T-t) - \ln \rho(t)}{\|\theta\|\sqrt{T-t}} \right) \right] \\ &\quad \times (\sigma^\top)^{-1} \theta, \\ X_1(t) &= \rho(t)^{-\frac{1}{\eta}\left(\frac{a}{\sigma_\rho}+1\right)} e^{\left[\left(\frac{1}{\eta}\left(\frac{a}{\sigma_\rho}+1\right)-1\right)r + \frac{1}{\eta}\left(\frac{a}{\sigma_\rho}+1\right)\left(\frac{1}{\eta}\left(\frac{a}{\sigma_\rho}+1\right)-1\right)\frac{\|\theta\|^2}{2}\right](T-t)} \\ &\quad \times \Phi \left(\frac{\ln c^* + \left(r + \left(\frac{1}{\eta}\left(\frac{a}{\sigma_\rho}+1\right) - \frac{1}{2}\right)\|\theta\|^2\right)(T-t) - \ln \rho(t)}{\|\theta\|\sqrt{T-t}} \right). \end{aligned}$$

Similarly, the portfolio replicating $\mathbf{1}_{\{\rho > c^*\}}$ is

$$\begin{aligned} \pi_2(t) &= \left[-\frac{1}{\|\theta\|\sqrt{T-t}\rho(t)} c^* \Phi' \left(\frac{\ln c^* + \left(r + \frac{\|\theta\|^2}{2}\right)(T-t) - \ln \rho(t)}{\|\theta\|\sqrt{T-t}} \right) \right] (\sigma^\top)^{-1} \theta, \\ X_2(t) &= e^{-r(T-t)} \left[1 - \Phi \left(\frac{\ln c^* + \left(r - \frac{1}{2}\|\theta\|^2\right)(T-t) - \ln \rho(t)}{\|\theta\|\sqrt{T-t}} \right) \right]. \end{aligned}$$

As a result, we derive (44).

On the other hand, we have

$$\begin{aligned} \frac{\partial \Gamma}{\partial \rho}(t, \rho) &= k_r(c^*)^{\frac{1}{\eta}\left(\frac{a}{\sigma_\rho}+1\right)} e^{\left(\left(\frac{1}{\eta}\left(\frac{a}{\sigma_\rho}+1\right)-1\right)r + \frac{1}{\eta}\left(\frac{a}{\sigma_\rho}+1\right)\left(\frac{1}{\eta}\left(\frac{a}{\sigma_\rho}+1\right)-1\right)\frac{\|\theta\|^2}{2}\right)(T-t)} \left(-\frac{1}{\eta} \left(\frac{a}{\sigma_\rho} + 1 \right) \right) \\ &\quad \times \rho^{-\frac{1}{\eta}\left(\frac{a}{\sigma_\rho}+1\right)-1} \Phi \left(\frac{\ln c^* + \left(r + \left(\frac{1}{\eta}\left(\frac{a}{\sigma_\rho}+1\right) - \frac{1}{2}\right)\|\theta\|^2\right)(T-t) - \ln \rho}{\|\theta\|\sqrt{T-t}} \right) < 0. \end{aligned}$$

Thus, $\Gamma(t, \rho)$ is strictly decreasing in ρ . Because

$$k_r e^{-r(T-t)} \left(1 - \Phi \left(\frac{\ln c^* + \left(r - \frac{1}{2}\|\theta\|^2\right)(T-t) - \ln \rho}{\|\theta\|\sqrt{T-t}} \right) \right)$$

is strictly increasing in ρ , we can deduct that $\Delta(t, \rho)$ is strictly decreasing in ρ . Finally, all of the

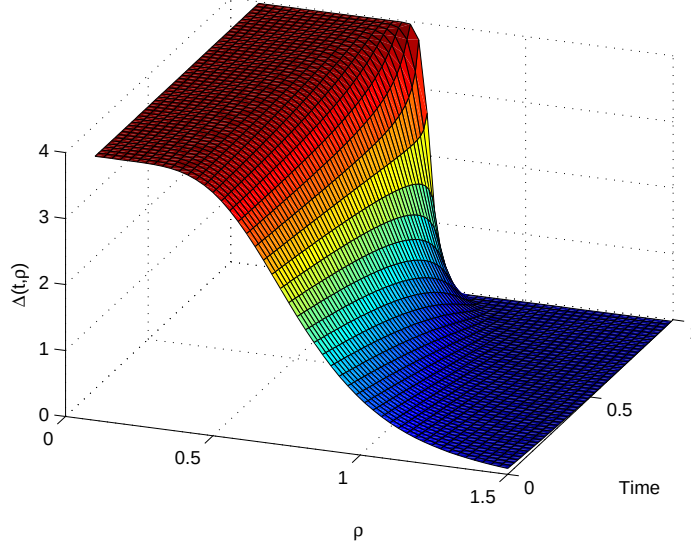


Figure 11: The quantity $\Delta(t, \rho)$ at different time t and pricing kernel value ρ .

limits in (45) can be derived easily. ■

From the monotonicity of $\Gamma(t, \rho)$ in ρ and (45), it follows that the optimal wealth process always lies above the portfolio insurance level, $e^{-r(T-t)}k_r$. On the other hand, $\Delta(t, \rho)$ is an interesting quantity. Recall that in classical expected utility maximization the optimal portfolio is given in the feedback form $\frac{1}{\eta}(\sigma^\top)^{-1}\theta X(t)$, i.e., the well-known Merton strategy. Thus, $\Delta(t, \rho)$ measures the deviation from the Merton strategy due to the presence of the probability weighting function. In view of (45), when ρ is very small (i.e., when the market is very good), $\Delta(t, \rho) \approx \frac{a}{\sigma_\rho} + 1 > 1$, indicating that the agent takes higher leverage than the Merton strategy does. Moreover, the magnitude of $\Delta(t, \rho)$ is (roughly) proportional to a , a parameter that measures the degree of hope. When ρ is very large (i.e., the market is very bad), $\Delta(t, \rho) \approx 0$, indicating that the agent takes little risky exposure due to fear and indicating the resulting portfolio insurance requirement. Figure 11 depicts $\Delta(t, \rho)$ with the parameters taking the values specified at the beginning of the present section.

7 Concluding Remarks

In this paper we have introduced and formulated a new portfolio choice model—the HF/A model—in continuous time. This model considers the role in decision making of three emotions: hope, fear and aspirations. Hope and fear are modeled through a reversed S-shaped probability weighting function from Quiggin’s RDU theory and aspirations are modeled by a probabilistic constraint from Lopes’ SP/A theory. Motivated by a comparative statics analysis on the optimal strategies in our model, these emotions have been further quantified via respective indices. In

particular, the hope and fear indices are related to the curvatures of the probability weighting function, whereas the level of aspirations can be measured by the degree to which the optimal payoff resembles that of a lottery ticket.

The SP/A theory motivated us to study the portfolio choice problem featuring hope, fear and aspirations. However, the ill-posedness result (Theorem 1) implies that the preference measure in SP/A theory is inappropriate in continuous-time portfolio choice modeling because of the linearity of the utility function. We have resolved this problem by introducing a concave utility function, a component taken from RDU theory. We have found, in our HF/A model, that both fear and the dislike of mean-preserving spreads lead to risk-averse behavior and that both hope and aspirations lead to risk-seeking behavior. However, we have shown that each of these four component plays its own distinct role in influencing investing behavior.

Appendix

A Proof of Theorem 5

The key step in the proof of Theorem 5 is to solve (16) for each $\lambda > 0$. First, we characterize the subset of feasible solutions in which the optimal solution to (16) must lie. Following the notation in Section 5.3, define

$$\mathbb{S}_\lambda^A := \{G(\cdot) \in \mathbb{G} \mid G(z) = b\mathbf{1}_{\{0 < z \leq z_0\}} + [b \vee \bar{G}_A(z)] \mathbf{1}_{\{z_0 < z < 1\}}, b \geq \bar{G}(z_0)\} \quad (46)$$

if $1 - \alpha \geq z_0$; and

$$\begin{aligned} \mathbb{S}_\lambda^A := \{G(\cdot) \in \mathbb{G} \mid G(z) = & b\mathbf{1}_{\{0 < z \leq 1-\alpha\}} + [b \vee A] \mathbf{1}_{\{1-\alpha < z \leq z_0\}} \\ & + [b \vee \bar{G}_A(z)] \mathbf{1}_{\{z_0 < z < 1\}}, b \geq \bar{G}(z_0)\} \end{aligned} \quad (47)$$

if $1 - \alpha \leq z_0$, where

$$\bar{G}_A(z) := \bar{G}(z)\mathbf{1}_{\{0 < z \leq 1-\alpha\}} + [\bar{G}(z) \vee A] \mathbf{1}_{\{1-\alpha < z < 1\}}, \quad 0 < z < 1$$

and $\bar{G}(\cdot)$ is defined as in (25).

Proposition 8 *Let Assumptions 5 and 7 hold. For any $G(\cdot) \in \mathbb{G}$, there exists a $\tilde{G}(\cdot) \in \mathbb{S}_\lambda$ such that $U_\lambda(\tilde{G}(\cdot)) \geq U_\lambda(G(\cdot))$ and the inequality becomes equality if and only if $G(\cdot) = \tilde{G}(\cdot)$.*

Proof First, consider the case in which $1 - \alpha \geq z_0$. For any feasible $G(\cdot)$, let $z_1 := \inf\{z \in (0, z_0] \mid G(z) > \bar{G}(z)\}$ with $\inf \emptyset = z_0$. Let $b = \bar{G}(z_1)$ and $z_2 := \inf\{z \in [z_0, 1) \mid \bar{G}(z) > b\}$.

Clearly, $z_0 \leq z_2 < 1$ and $\bar{G}(z_2) = \bar{G}(z_1) = b$. If $z_2 \leq 1 - \alpha$, define

$$\tilde{G}(z) := \begin{cases} b & 0 < z \leq z_2 \\ \bar{G}(z) & z_2 < z \leq 1 - \alpha \\ \bar{G}_A(z) & 1 - \alpha < z < 1. \end{cases}$$

If $z_2 > 1 - \alpha$, define

$$\tilde{G}(z) := \begin{cases} b & 0 < z \leq 1 - \alpha \\ b \vee \bar{G}_A(z) & 1 - \alpha < z < 1. \end{cases}$$

In both cases, we have $\tilde{G}(\cdot) \in \mathbb{S}_\lambda^A$.

If $z_2 \leq 1 - \alpha$, we have

$$\begin{aligned} U_\lambda(G(\cdot)) &= \int_0^1 f(G(z), z) dz \\ &= \int_0^{z_1} f(G(z), z) dz + \int_{z_1}^{z_2} f(G(z), z) dz \\ &\quad + \int_{z_2}^{1-\alpha} f(G(z), z) dz + \int_{1-\alpha}^1 f(G(z), z) dz \\ &\leq \int_0^{z_1} f(\bar{G}(z_1), z) dz + \int_{z_1}^{z_2} f(\bar{G}(z_1), z) dz \\ &\quad + \int_{z_2}^{1-\alpha} f(\bar{G}(z), z) dz + \int_{1-\alpha}^1 f(\bar{G}_A(z), z) dz \\ &= \int_0^1 f(\tilde{G}(z), z) dz, \end{aligned}$$

and the inequality becomes equality if and only if $G(\cdot) = \tilde{G}(\cdot)$. If $z_2 > 1 - \alpha$, we have

$$\begin{aligned} U_\lambda(G(\cdot)) &= \int_0^1 f(G(z), z) dz \\ &= \int_0^{z_1} f(G(z), z) dz + \int_{z_1}^{1-\alpha} f(G(z), z) dz + \int_{1-\alpha}^1 f(G(z), z) dz \\ &\leq \int_0^{z_1} f(\bar{G}(z_1), z) dz + \int_{z_1}^{1-\alpha} f(\bar{G}(z_1), z) dz + \int_{1-\alpha}^1 f(b \vee \bar{G}_A(z), z) dz \\ &= \int_0^1 f(\tilde{G}(z), z) dz, \end{aligned}$$

and the inequality becomes equality if and only if $G(\cdot) = \tilde{G}(\cdot)$.

Next, we consider the case in which $1 - \alpha < z_0$. For any feasible $G(\cdot)$, again let $z_1 := \inf\{z \in (0, z_0] \mid G(z) > \bar{G}(z)\}$ with $\inf \emptyset = z_0$. Let $b = \bar{G}(z_1)$. If $b \geq A$, define $z_2 := \inf\{z \in [z_0, 1) \mid$

$\bar{G}(z) > b\}$, and

$$\tilde{G}(z) := \begin{cases} b & 0 < z \leq z_2 \\ \bar{G}(z) & z_2 < z < 1. \end{cases}$$

If $b < A$, then the feasibility of $G(\cdot)$ implies $z_1 \leq 1 - \alpha$. Define

$$\tilde{G}(z) := \begin{cases} b & 0 < z \leq 1 - \alpha \\ \bar{G}_A(z) & 1 - \alpha < z < 1. \end{cases}$$

In both cases, $\tilde{G}(z) \in \mathbb{S}_\lambda^A$. If $b \geq A$, we have

$$\begin{aligned} U_\lambda(G(\cdot)) &= \int_0^1 f(G(z), z) dz \\ &= \int_0^{z_1} f(G(z), z) dz + \int_{z_1}^{z_2} f(G(z), z) dz + \int_{z_2}^1 f(G(z), z) dz \\ &\geq \int_0^{z_1} f(\bar{G}(z_1), z) dz + \int_{z_1}^{z_2} f(\bar{G}(z_1), z) dz + \int_{z_2}^1 f(\bar{G}(z), z) dz \\ &= U_\lambda(\tilde{G}(\cdot)), \end{aligned}$$

and the inequality becomes equality if and only if $G(\cdot) = \tilde{G}(\cdot)$. If $b < A$, we have

$$\begin{aligned} U_\lambda(G(\cdot)) &= \int_0^1 f(G(z), z) dz \\ &= \int_0^{z_1} f(G(z), z) dz + \int_{z_1}^{1-\alpha} f(G(z), z) dz + \int_{1-\alpha}^1 f(G(z), z) dz \\ &\geq \int_0^{z_1} f(\bar{G}(z_1), z) dz + \int_{z_1}^{1-\alpha} f(\bar{G}(z_1), z) dz + \int_{z_2}^1 f(\bar{G}_A(z), z) dz \\ &= U_\lambda(\tilde{G}(\cdot)), \end{aligned}$$

and the inequality becomes equality if and only if $G(\cdot) = \tilde{G}(\cdot)$. ■

In view of Proposition 8, we need only consider the following problem:

$$\begin{aligned} \text{Max}_{G(\cdot)} \quad & U_\lambda(G(\cdot)) = \int_0^1 [u(G(z))w'(1-z) - \lambda G(z)F_\rho^{-1}(1-z)] dz \\ \text{subject to} \quad & G(\cdot) \in \mathbb{S}_\lambda^A, \quad G((1-\alpha)+) \geq A, \quad G(0+) \geq 0, \end{aligned} \tag{48}$$

which is an optimization problem over the real line. The following proposition provides the result.

Before we state the proposition and its proof, however, let us recall $\varphi(\cdot)$ in (27) and $\psi(\cdot)$ in (31) and their properties. $\varphi(\cdot)$ is strictly increasing on $(0, z_0)$ and strictly decreasing on $(z_0, 1)$. Furthermore, $\varphi(\cdot)$ is positive on $(0, z^*)$ and negative on $(z^*, 1)$. $\psi(\cdot)$ is strictly decreasing on $(0, z^*)$

and strictly increasing on $(z^*, 1)$. In addition, $\psi(z) > M(z)$ on $(0, z^*)$, $\psi(z) < M(z)$ on $(z^*, 1)$ and $\psi(z^*) = M(z^*)$.

Proposition 9 *Let Assumptions 5-7 hold. Let z^* be the unique root of (27) and recall $\psi(\cdot)$ in (31).*

1. *If $0 < 1 - \alpha < z^*$, then $\psi(1 - \alpha) > \psi(z^*) = M(z^*)$, and the unique optimal solution to (16) is given as*

$$G_\lambda^*(z) = \begin{cases} \bar{G}(z^*)\mathbf{1}_{\{0 < z \leq z^*\}} + \bar{G}(z)\mathbf{1}_{\{z^* < z \leq 1\}} & \frac{\lambda}{w'(A)} \leq M(z^*), \\ A\mathbf{1}_{\{0 < z \leq z_0\}} + \bar{G}_A(z)\mathbf{1}_{\{z_0 < z \leq 1\}} & M(z^*) \leq \frac{\lambda}{w'(A)} \leq \psi(1 - \alpha), \\ (u')^{-1}\left(\frac{\lambda}{\psi(1 - \alpha)}\right)\mathbf{1}_{\{0 < z \leq 1 - \alpha\}} + \bar{G}_A(z)\mathbf{1}_{\{1 - \alpha < z \leq 1\}} & \frac{\lambda}{w'(A)} \geq \psi(1 - \alpha). \end{cases} \quad (49)$$

2. *If $z^* \leq 1 - \alpha < 1$, then the unique optimal solution to (16) is given as*

$$G_\lambda^*(z) = \bar{G}(z^*)\mathbf{1}_{\{0 < z \leq z^*\}} + \bar{G}_A(z)\mathbf{1}_{\{z^* < z \leq 1\}}. \quad (50)$$

Proof The proof is split into three cases: (i) $1 - \alpha \geq z^*$; (ii) $z_0 < 1 - \alpha < z^*$; and (iii) $1 - \alpha < z_0$.

- (i) First, consider the case in which $1 - \alpha \geq z^*$. If $A \leq \bar{G}(1 - \alpha)$, i.e., $\frac{\lambda}{w'(A)} \leq M(1 - \alpha)$, then the optimal solution to (16) when $A = 0$ in Proposition 5 automatically satisfies the additional aspiration constraint $G((1 - \alpha) +) \geq A$, and therefore it is also optimal in the presence of A . It is easy to check that the optimal solution in (28) coincides with the one in (50) when $A \leq \bar{G}(1 - \alpha)$. Thus, we can assume that $A \geq \bar{G}(1 - \alpha)$, i.e., $\frac{\lambda}{w'(A)} \geq M(1 - \alpha)$. Define

$$V(b) := U(b\mathbf{1}_{\{0 < z \leq z_0\}} + [b \vee \bar{G}_A(z)]\mathbf{1}_{\{z_0 < z < 1\}}), \quad b \geq \bar{G}(z_0).$$

Consider $V(\cdot)$ in three different domains. When $\bar{G}(z_0) \leq b \leq \bar{G}(1 - \alpha)$, let $y := \bar{G}^{-1}(b)$, where $\bar{G}^{-1}(\cdot)$ is the inverse function of $\bar{G}(\cdot)$ when restricted on $[z_0, 1)$. We have

$$\begin{aligned} V(b) &= u(b) \int_0^y w'(1 - z)dz - \lambda b \int_0^y F_\rho^{-1}(1 - z)dz \\ &\quad + \int_y^1 u(\bar{G}_A(z))w'(1 - z)dz - \lambda \int_y^1 \bar{G}_A(z)F_\rho^{-1}(1 - z)dz. \end{aligned}$$

Clearly,

$$\begin{aligned} V'(b) &= u'(b) \int_0^y w'(1-z)dz - \lambda \int_0^y F_\rho^{-1}(1-z)dz \\ &= \frac{\lambda}{M(y)} \varphi(y). \end{aligned}$$

Because $\bar{G}(z_0) \leq b \leq \bar{G}(1-\alpha)$, $z_0 \leq y \leq 1-\alpha$. Recalling that $z_0 < z^* \leq 1-\alpha$, we conclude that $b_1^* := \bar{G}(z^*)$ uniquely maximizes $V(\cdot)$ in the interval $[\bar{G}(z_0), \bar{G}(1-\alpha)]$.

When $\bar{G}(1-\alpha) < b < A$, we have

$$\begin{aligned} V(b) &= u(b) \int_0^{1-\alpha} w'(1-z)dz - \lambda b \int_0^{1-\alpha} F_\rho^{-1}(1-z)dz \\ &\quad + \int_{1-\alpha}^1 u(\bar{G}_A(z))w'(1-z)dz - \lambda \int_{1-\alpha}^1 \bar{G}_A(z)F_\rho^{-1}(1-z)dz, \end{aligned}$$

and

$$\begin{aligned} V'(b) &= u'(b) \int_0^{1-\alpha} w'(1-z)dz - \lambda \int_0^{1-\alpha} F_\rho^{-1}(1-z)dz \\ &< \lambda \left[\frac{1}{M(1-\alpha)} \int_0^{1-\alpha} w'(1-z)dz - \int_0^{1-\alpha} F_\rho^{-1}(1-z)dz \right] \leq 0, \end{aligned}$$

where the first inequality is the case because $b > \bar{G}(1-\alpha)$ and the last inequality is the case because $1-\alpha \geq z^*$ and $\varphi(z) < 0$ on $(z^*, 1)$. Therefore, $V(\cdot)$ is strictly decreasing in the interval $[\bar{G}(1-\alpha), A]$.

When $b \geq A$, let $y := \bar{G}^{-1}(b) \geq 1-\alpha$. Then, we have

$$\begin{aligned} V(b) &= u(b) \int_0^y w'(1-z)dz - \lambda b \int_0^y F_\rho^{-1}(1-z)dz \\ &\quad + \int_y^1 u(\bar{G}_A(z))w'(1-z)dz - \lambda \int_y^1 (\bar{G}_A(z))F_\rho^{-1}(1-z)dz, \end{aligned}$$

and

$$\begin{aligned} V'(b) &= u'(b) \int_0^y w'(1-z)dz - \lambda \int_0^y F_\rho^{-1}(1-z)dz \\ &= \frac{\lambda}{M(y)} \varphi(y) \leq 0, \end{aligned}$$

where the last inequality is due to $z^* \leq 1-\alpha$ and $\varphi(z) < 0$ on $(z^*, 1)$. Therefore, $V(\cdot)$ is decreasing on $(A, +\infty)$.

To summarize, we conclude that $V(\cdot)$ obtains its maximum value at $b_1^* = \bar{G}(z^*)$, and therefore the optimal solution is given in (50).

- (ii) Next, consider the case in which $z_0 < 1 - \alpha < z^*$. Using the same argument as for case (i), we can assume that $A \geq \bar{G}(1 - \alpha)$, i.e., $\frac{\lambda}{w'(A)} \geq M(1 - \alpha)$. Because $1 - \alpha < z^*$, we have $M(1 - \alpha) < \psi(1 - \alpha)$. Again, define

$$V(b) := U(b \mathbf{1}_{\{0 < z \leq z_0\}} + [b \vee \bar{G}_A(z)] \mathbf{1}_{\{z_0 < z < 1\}}), \quad b \geq \bar{G}(z_0).$$

When $\bar{G}(z_0) \leq b \leq \bar{G}(1 - \alpha)$, letting $y := \bar{G}^{-1}(b) \leq 1 - \alpha$, we have

$$\begin{aligned} V(b) &= u(b) \int_0^y w'(1 - z) dz - \lambda b \int_0^y F_\rho^{-1}(1 - z) dz \\ &\quad + \int_y^1 u(\bar{G}_A(z)) w'(1 - z) dz - \lambda \int_y^1 (\bar{G}_A(z)) F_\rho^{-1}(1 - z) dz. \end{aligned}$$

Clearly,

$$\begin{aligned} V'(b) &= u'(b) \int_0^y w'(1 - z) dz - \lambda \int_0^y F_\rho^{-1}(1 - z) dz \\ &= \frac{\lambda}{M(y)} \varphi(y) > 0, \end{aligned}$$

where the last inequality is the case because $1 - \alpha < z^*$ and $\varphi(z) > 0$ on $(0, z^*)$. Therefore, $V(\cdot)$ is strictly increasing in the interval $[\bar{G}(z_0), \bar{G}(1 - \alpha)]$.

When $\bar{G}(1 - \alpha) < b < A$, we have

$$\begin{aligned} V(b) &= u(b) \int_0^{1-\alpha} w'(1 - z) dz - \lambda b \int_0^{1-\alpha} F_\rho^{-1}(1 - z) dz \\ &\quad + \int_{1-\alpha}^1 u(\bar{G}_A(z)) w'(1 - z) dz - \lambda \int_{1-\alpha}^1 \bar{G}_A(z) F_\rho^{-1}(1 - z) dz, \end{aligned}$$

and

$$V'(b) = u'(b) \int_0^{1-\alpha} w'(1 - z) dz - \lambda \int_0^{1-\alpha} F_\rho^{-1}(1 - z) dz.$$

If $M(1 - \alpha) \leq \frac{\lambda}{w'(A)} \leq \psi(1 - \alpha)$, we have

$$\begin{aligned} V'(b) &> u'(A) \int_0^{1-\alpha} w'(1 - z) dz - \lambda \int_0^{1-\alpha} F_\rho^{-1}(1 - z) dz \\ &= u'(A) \int_0^{1-\alpha} F_\rho^{-1}(1 - z) dz \left[\psi(1 - \alpha) - \frac{\lambda}{u'(A)} \right] \geq 0, \end{aligned}$$

and, consequently, $V(\cdot)$ is strictly increasing in the interval $[\bar{G}(1 - \alpha), A]$. If $\frac{\lambda}{u'(A)} \geq \psi(1 - \alpha)$, then it is easy to see that the unique optimizer of $V(\cdot)$ in this interval is $b_2^* := (u')^{-1} \left(\frac{\lambda}{\psi(1 - \alpha)} \right)$.

When $b \geq A$, let $y := \bar{G}^{-1}(b)$. Then, we have

$$\begin{aligned} V(b) &= u(b) \int_0^y w'(1-z)dz - \lambda b \int_0^y F_\rho^{-1}(1-z)dz \\ &\quad + \int_y^1 u(\bar{G}_A(z))w'(1-z)dz - \lambda \int_y^1 \bar{G}_A(z)F_\rho^{-1}(1-z)dz, \end{aligned}$$

and

$$\begin{aligned} V'(b) &= u'(b) \int_0^y w'(1-z)dz - \lambda \int_0^y F_\rho^{-1}(1-z)dz \\ &= \frac{\lambda}{M(y)}\varphi(y). \end{aligned}$$

If $\frac{\lambda}{u'(A)} \geq M(z^*)$, then $\bar{G}^{-1}(A) \geq z^*$. Consequently, $y > \bar{G}^{-1}(A) \geq z^*$ and $V'(b) < 0$ because $\varphi(\cdot)$ is negative on $(z^*, 1)$. In other words, $V(\cdot)$ is strictly decreasing on this interval. If $\frac{\lambda}{u'(A)} \leq M(z^*)$, then $\bar{G}^{-1}(A) \leq z^*$ and, consequently, the unique maximizer $b_3^* = (u')^{-1}\left(\frac{\lambda}{M(z^*)}\right)$.

Finally, because $z_0 \leq 1 - \alpha < z^*$, $M(1 - \alpha) < M(z^*) = \psi(z^*) < \psi(1 - \alpha)$. Then, we can sum up the results and conclude that the optimizer of $V(\cdot)$ is $(u')^{-1}\left(\frac{\lambda}{\psi(1 - \alpha)}\right)$ if $\frac{\lambda}{u'(A)} \geq \psi(1 - \alpha)$, is A if $M(z^*) \leq \frac{\lambda}{u'(A)} \leq \psi(1 - \alpha)$ and is $(u')^{-1}\left(\frac{\lambda}{M(z^*)}\right)$ if $M(1 - \alpha) \leq \frac{\lambda}{u'(A)} \leq M(z^*)$. Therefore, the optimal solution is given in (49).

- (iii) Finally, consider the case in which $1 - \alpha \leq z_0$. Using the same arguments as in (i) and (ii), we can assume that $A \geq \bar{G}(z_0)$, i.e., $\frac{\lambda}{u'(A)} \geq M(z_0)$. Again, we let

$$V(b) := U(b\mathbf{1}_{\{0 < z \leq 1 - \alpha\}} + [b \vee A]\mathbf{1}_{\{1 - \alpha < z \leq z_0\}} + [b \vee \bar{G}_A(z)]\mathbf{1}_{\{z_0 < z < 1\}}), \quad b \geq M(z_0).$$

We first optimize $V(\cdot)$ in $[A, +\infty)$. In this interval, we have

$$V(b) = U(b\mathbf{1}_{\{0 < z \leq z_0\}} + [b \vee \bar{G}(z)]\mathbf{1}_{\{z_0 < z < 1\}}),$$

which is the same as the function $V(\cdot)$ in the proof of Proposition 5 after making the transformation $y = \bar{G}^{-1}(b)$. From the proof of Proposition 5, we can deduce that the maximizer of $V(\cdot)$ on $[A, +\infty)$ is $\bar{G}(z^*)$ if $\frac{\lambda}{u'(A)} \leq M(z^*)$ and is A if $\frac{\lambda}{u'(A)} \geq M(z^*)$.

Next, we consider $V(b)$ for other possible b . If $A \leq \bar{G}(1 - \alpha)$, i.e., $\frac{\lambda}{u'(A)} \leq M(1 - \alpha)$, then it is easy to see that $V(\cdot)$ is strictly increasing on $[\bar{G}(z_0), A]$. If $A \geq \bar{G}(1 - \alpha)$, i.e., $\frac{\lambda}{u'(A)} \geq M(1 - \alpha)$, then it is easy to check that $V(\cdot)$ is strictly increasing in $[\bar{G}(z_0), \bar{G}(1 - \alpha)]$. Thus, we need only consider the case in which $A > \bar{G}(1 - \alpha)$, i.e., $\frac{\lambda}{u'(A)} > M(1 - \alpha)$, and

check $V(b)$ for $\bar{G}(1 - \alpha) < b < A$. In this case, we have

$$\begin{aligned} V(b) &= u(b) \int_0^{1-\alpha} w'(1-z)dz - \lambda b \int_0^{1-\alpha} F_\rho^{-1}(1-z)dz \\ &\quad + \int_{1-\alpha}^1 u(\bar{G}_A(z))w'(1-z)dz - \lambda \int_{1-\alpha}^1 \bar{G}_A(z)F_\rho^{-1}(1-z)dz, \end{aligned}$$

and

$$V'(b) = u'(b) \int_0^{1-\alpha} w'(1-z)dz - \lambda \int_0^{1-\alpha} F_\rho^{-1}(1-z)dz.$$

Because $1 - \alpha \leq z_0 < z^*$, we have $\psi(1 - \alpha) > M(1 - \alpha)$. Now, when $M(1 - \alpha) \leq \frac{\lambda}{u'(A)} \leq \psi(1 - \alpha)$,

$$V'(b) > u'(A) \left(\int_0^{1-\alpha} F_\rho^{-1}(1-z)dz \right) \left[\psi(1 - \alpha) - \frac{\lambda}{u'(A)} \right] \geq 0,$$

where the first inequality is due to $b < A$. Consequently, $V(\cdot)$ is strictly increasing in $[\bar{G}(1 - \alpha), A]$. When $\frac{\lambda}{u'(A)} \geq \psi(1 - \alpha)$, it is easy to see that the optimizer of $V(\cdot)$ in this interval is $(u')^{-1} \left(\frac{\lambda}{\psi(1-\alpha)} \right)$.

To summarize, the unique optimizer of $V(\cdot)$ is $\bar{G}(z^*)$ if $\frac{\lambda}{u'(A)} \leq M(z^*)$, is A if $M(z^*) \leq \frac{\lambda}{u'(A)} \leq \psi(1 - \alpha)$ and is $(u')^{-1} \left(\frac{\lambda}{\psi(1-\alpha)} \right)$ if $\frac{\lambda}{u'(A)} \geq \psi(1 - \alpha)$. Therefore, the optimal solution is given in (49).

Ultimately, the uniqueness of the optimal solution can be easily derived from the above proof. ■

Proof of Theorem 5 Let

$$\mathcal{X}(\lambda) = \int_0^1 G_\lambda^*(z)F_\rho^{-1}(1-z)dz, \quad \lambda > 0$$

where $G_\lambda^*(\cdot)$ is as given in (49) or (50), depending on the value of α . By Assumption 6, $\mathcal{X}(\cdot)$ is finite and decreasing on $(0, +\infty)$. Furthermore, because ρ is atomless, by applying the monotone convergence theorem, we conclude that $\mathcal{X}(\cdot)$ is continuous and

$$\lim_{\lambda \uparrow +\infty} \mathcal{X}(\lambda) = AE \left[\rho \mathbf{1}_{\{\rho \leq F_\rho^{-1}(\alpha)\}} \right], \quad \lim_{\lambda \downarrow 0} \mathcal{X}(\lambda) = +\infty.$$

Therefore, we can find λ^* such that $\mathcal{X}(\lambda^*) = x_0$ and, consequently, $G_{\lambda^*}^*(\cdot)$ is optimal to (14). Noticing that $\mathcal{X}(u'(A)M(z^*)) = x_r$ and that $\mathcal{X}(u'(A)\psi(1 - \alpha)) = x_p$, the optimal solution to (13), $X^* := G_{\lambda^*}^*(1 - F_\rho(\rho))$, is exactly as given in (34)-(37). ■

References

- ABDELLAOUI, M. (2000): “Parameter-Free Elicitation of Utility and Probability Weighting Functions,” *Management Science*, 46(11), 1497–1512.
- (2002): “A Genuine Rank-Dependent Generalization of The Von Neumann-Morgenstern Expected Utility Theorem,” *Econometrica*, 70(1), 717–736.
- ABDELLAOUI, M., H. BLEICHRODT, AND C. PARASCHIV (2007): “Loss Aversion Under Prospect Theory: A Parameter-Free Measurement,” *Management Science*, 53(10), 1659–1674.
- BARBERIS, N. (2012): “A Model of Casino Gambling,” *Management Science*, 58(1), 35–51.
- BASAK, S. (1995): “A General Equilibrium Model of Portfolio Insurance,” *Review of Financial Studies*, 8(4), 1059–1090.
- BENARTZI, S., AND R. H. THALER (1995): “Myopic Loss Aversion and the Equity Premium Puzzle,” *Quarterly Journal of Economics*, 110(1), 73–92.
- BJÖRK, T., A. MURGOI, AND X. Y. ZHOU (2011): “Mean-Variance Portfolio Optimization with State Dependent Risk Aversion,” To appear in *Mathematical Finance*.
- BLEICHRODT, H., AND J. L. PINTO (2000): “A Parameter-Free Elicitation of the Probability Weighting Function in Medical Decision Analysis,” *Management Science*, 46(11), 1485–1496.
- CAMERER, C. F., AND T.-H. HO (1994): “Violations of the betweenness axiom and nonlinearity in probability,” *Journal of Risk and Uncertainty*, 8(2), 167–196.
- CARLIER, G., AND R.-A. DANA (2006): “Law invariant concave utility functions and optimization problems with monotonicity and comonotonicity constraints,” *Statistics and Decisions*, 24(1), 127–152.
- (2011): “Optimal Demand for Contingent Claims When Agents Have Law Invariant Utilities,” *Mathematical Finance*, 21(2), 169–201.
- EKELAND, I., AND A. LAZRAK (2006): “Being serious about non-commitment: subgame perfect equilibrium in continuous time,” Working Paper.
- EKELAND, I., AND T. PIRVU (2008): “Investment and consumption without commitment,” *Mathematics and Financial Economics*, 2(1), 57–86.
- GROSSMAN, S. J., AND Z. ZHOU (1996): “Equilibrium Analysis of Portfolio Insurance,” *Journal of Finance*, 51(4), 1379–1403.
- HE, X. D., AND X. Y. ZHOU (2011a): “Portfolio Choice under Cumulative Prospect Theory: An Analytical Treatment,” *Management Science*, 57(2), 315–331.

- (2011b): “Portfolio Choice via Quantiles,” *Mathematical Finance*, 21(2), 203–231.
- INGERSOLL, J. E. (2008): “Non-Monotonicity of the Tversky-Kahneman Probability-Weighting Function: A Cautionary Note,” *European Financial Management*, 14(3), 385–390.
- (2011): “Cumulative Prospect Theory and the Representative Investor,” Working Paper.
- JIN, H. Q., Z. Q. XU, AND X. Y. ZHOU (2007): “A Convex Stochastic Optimization Problem Arising from Portfolio Selection,” *Mathematical Finance*, 18(1), 171–183.
- JIN, H. Q., AND X. Y. ZHOU (2008): “Behavioral Portfolio Selection in Continuous Time,” *Mathematical Finance*, 18(3), 385–426.
- KAHNEMAN, D., AND A. TVERSKY (1979): “Prospect Theory: An Analysis of Decision Under Risk,” *Econometrica*, 47(2), 263–291.
- KARATZAS, I., AND S. E. SHREVE (1998): *Methods of Mathematical Finance*. Springer, New York.
- KRAMKOV, D., AND W. SCHACHERMAYER (1999): “The Asymptotic Elasticity of Utility Functions and Optimal Investment in Incomplete Markets,” *Annals of Applied Probability*, 9(3), 904–950.
- LOPES, L. L. (1987): “Between Hope and Fear: The Psychology of Risk,” *Advances in Experimental Social Psychology*, 20, 255–295.
- LOPES, L. L., AND G. C. ODEN (1999): “The Role of Aspiration Level in Risk Choice: A Comparison of Cumulative Prospect Theory and SP/A Theory,” *Journal of Mathematical Psychology*, 43(2), 286–313.
- LUCAS, D. J. (1994): “Asset pricing with undiversifiable income risk and short sales constraints: Deepening the equity premium puzzle,” *Review of Financial Studies*, 34(3), 325–341.
- MACHINA, M. J. (1982): ““Expected Utility” Analysis without the Independence Axiom,” *Econometrica*, 50(2), 277–322.
- MEHRA, R., AND E. C. PRESCOTT (1985): “The Equity Premium: A Puzzle,” *Journal of Monetary Economics*, 15(2), 145–161.
- PRELEC, D. (1998): “The Probability Weighting Function,” *Econometrica*, 66(3), 497–527.
- QUIGGIN, J. (1982): “A Theory of Anticipated Utility,” *Journal of Economic Behavior and Organization*, 3, 323–343.
- SCHMEIDLER, D. (1989): “Subjective Probability and Expected Utility without Additivity,” *Econometrica*, 57(3), 571–587.

- SHEFRIN, H., AND M. STATMAN (2000): “Behavioral Portfolio Theory,” *Journal of Financial and Quantitative Analysis*, 35(2), 127–151.
- STARMER, C. (2000): “Developments in Non-Expected Utility Theory: The Hunt for a Descriptive Theory of Choice under Risk,” *Journal of Economic Literature*, 38(2), 332–382.
- TVERSKY, A., AND C. R. FOX (1995): “Weighing Risk and Uncertainty,” *Psychological Review*, 102(2), 269–283.
- TVERSKY, A., AND D. KAHNEMAN (1992): “Advances in Prospect Theory: Cumulative Representation of Uncertainty,” *Journal of Risk and Uncertainty*, 5(4), 297–323.
- WANG, S. S. (2000): “A Class of Distortion Operators for Pricing Financial and Insurance Risks,” *Journal of Risk and Insurance*, 67(1), 15–36.
- WU, G., AND R. GONZALEZ (1996): “Curvature of the Probability Weighting Function,” *Management Science*, 42(12), 1676–1690.
- YAARI, M. E. (1987): “The Dual Theory of Choice under Risk,” *Econometrica*, 55(1), 95–115.