

Ruin probability

Dan Kucеровsky¹

Department of Mathematics and Statistics University of New Brunswick Fredericton, New Brunswick,
Canada

Abstract

We discuss studying ruin probability by Levy processes, and in particular the problem of approximating them through compound sum of exponentials leads us to some interesting mathematical problems. Complex analysis and Weiner-Hopf factorization turn out to be useful tools.

Email: dkucеров@unb.ca

¹speaker

On the importance of data mining methods in actuarial science

Shima Ara¹,

Saman Insurance Company, Tehran, Iran

Abstract

Nowadays, by rapid advancements in technology, the actuaries often encounter large amount of observations from various insurance fields. Modeling and analyzing such big data raise several challenges in the application. In recent years, data mining methods attracted many interest in various fields, as well as insurance. Actuaries should be ready to combine their traditional knowledge with the data mining approach to deal with the new world of big data. In this talk, we will review the impact of these methods in actuarial science and present findings of employing such methods on life and nonlife insurance in Saman insurance company.

¹speaker

Mortality modeling of skin cancer patients with actuarial applications

Amin Hassan Zadeh¹

Shahid Beheshti University, Tehran, Iran

Raoufeh Asghari

Shahid Beheshti University, Tehran, Iran

Abstract

In this article, the Markovian aging process is used to model mortality of patients with skin cancer. The time till death is assumed to have a phase-type distribution (which is defined in a Markov chain environment) with interpretable parameters. The underlying continuous-time Markov chain has one absorbing state (death) and $n_x + 1$ (x is the age when the patient is diagnosed with cancer) transient states. Each transient state represents a physiological age, and aging is transitions from one physiological age to the next one until the process reaches to its end. The transition can occur from any other state to the absorbing state. For patients with skin cancer in the United States, we estimate unknown parameters related to the aging process that can be useful for comparing the physiological age processes of patients with cancer and healthy people. For different age intervals, we estimate physiological age parameters for both males and females. The index of conditional expected physiological age of the patients with skin cancer at given ages are calculated and compared with the US total population. By using the bootstrap techniques, confidence bands and confidence intervals are constructed for the estimated survival curves and aging process parameters, respectively. The fitting results have been used for pricing the substandard annuities.

Keywords:

Aging process; mortality modeling; phase-type distribution; skin cancer; substandard annuities, bootstrap.

¹speaker

Use of Decision Tree in Classifying Customers of Third-Party Liability Motor Insurance

Farbod Khanizadeh¹

Insurance Research Center, Tehran, Iran

Mohammadreza Asghari Oskoei

Allameh Tabataba'i University, Tehran, Iran

Azadeh Bahador

Insurance Research Center, Tehran, Iran

Abstract

The purpose of this study was to present a pattern through which insurance companies are enabled to classify the risk level of their customer and to predict the possibility of future claims. We have analyzed an insurance claim dataset provided by an Iranian insurance company with a sample size of 6000. According to the structure of the dataset, a supervised learning algorithm was used to describe the underlying relationships between variables. The proposed algorithm, decision tree, was implemented using Python programming language. Based on the results, age, vehicle type and marital status were the main three factors contributing in prediction of claims.

Keywords: Decision Tree, Supervised Learning, Machine Learning, Classification, TPL.

1 Introduction

Motor third party liability insurance (TPL) is a type of policy financially protecting third parties against both physical damage and bodily injury caused by car accidents. Under third party insurance Act, TPL is a compulsory insurance product in Iran. This is a sound reason for insurers to make an attempt to have their customers' risks assessed as accurately as possible. Furthermore, Iran insurance industry is going to experience the transition from auto-based TPL to the driver-based one [1]. Currently premiums are calculated solely according to auto features. By driver-based policies we mean the situation in which both drivers' and autos' characteristics would affect premium rates.

The new regulations are going into effect by the end of 2021. Therefore, insurance companies need to review their risk assessment methods and make necessary amendments. In fact, in near future drivers' traits will also play a crucial role in TPL ratemaking. Hence in this article we worked with a dataset of features describing both drivers and vehicles. Basically we built a model to classify customers based on their risk factors divided into demographic information and some significant properties of cars ([2], [3], [4], [5]).

Decision tree is one of the common supervised learning algorithms used for both classification and regression. The method can be applied across a broad range of disciplines ([6], [7], [8], [9], [10], [11]). Due to the nature of insurance industry, risk classification, decision tree has received considerable attention within this field ([12], [13],

¹ speaker

[14], [15]). The focus of this article was to develop a tree classifier to help insurers predict their high and low risk policyholders ([16], [17], [18]).

Having gathered the data set, four steps were followed to build up the model:

1. Data preprocessing (replacing missing values, encoding categorical variables, split data into training and test set)
2. Build the tree classifier
3. Cross-validation and model evaluation
4. Avoid over-fitting (pre-pruning)
5. Model visualization

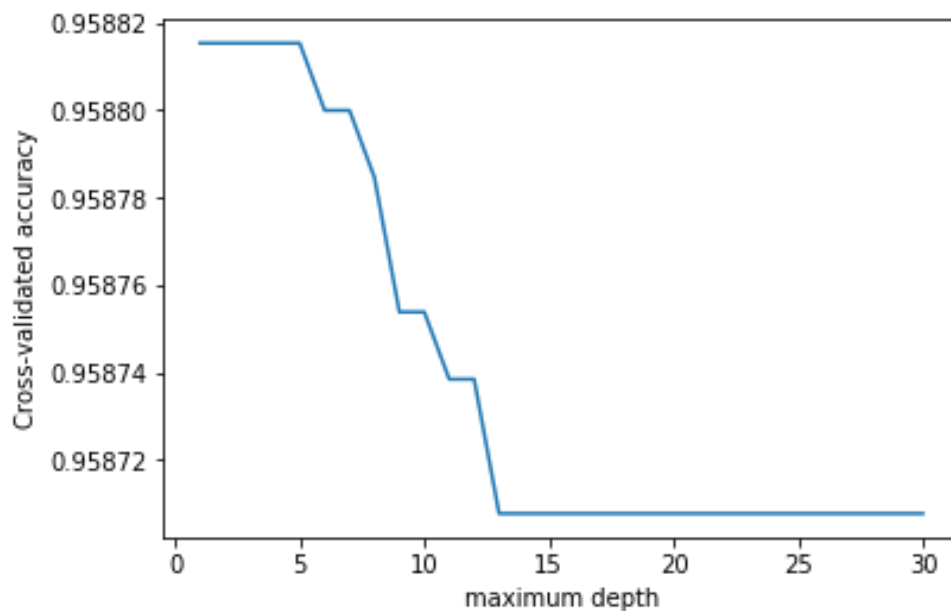
Regarding step two, we used the "Gini Index" as a measure of node impurity. Basically in building the binary tree, Gini index is the splitting criterion at each node and would determine the importance of each feature. Gini index is defined as:

$$\text{Gini} = 1 - \sum (P_i)^2$$

Where P_i corresponds to the probability of each class. The next section would present key findings and outcomes of the research. Note that the learning algorithm was evaluated on a dataset containing information regarding TLP claims. In the given dataset there were 6000 observations for 5 features namely; policy holder age, type of vehicle, marital status, gender and the dependent variable claim with two levels (Yes/No).

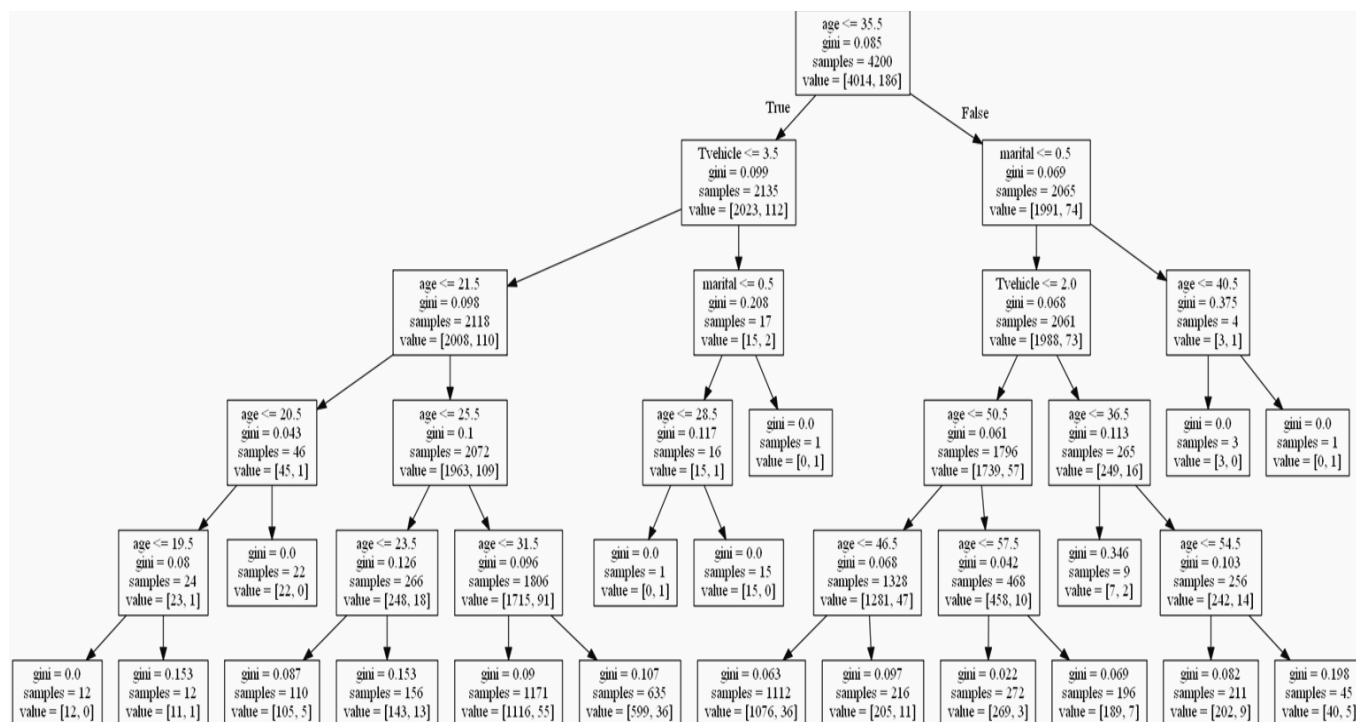
2 Main results

To obtain a reliable model we evaluated the accuracy of our classifier over both training and test sets. This was conducted through specifying the optimum value of a hyper parameter maximum depth of the tree. With the help of k-fold cross validation (for $k=10$) we tried thirty values for the depth and following figure was derived:



As shown in the figure, accuracy reaches lower score as we increase the maximum depth of tree. The highest accuracy belongs to the first five values. Therefore, we pick the highest depth corresponds to the best accuracy score (i.e. max_depth=5).

According to the above figure we have prevented a model from over-fitting via pre-pruning method. In fact the hyper parameter max_depth is used to set the maximum depth of our classifier. The following diagram depicts the decision tree of height 5:



The model consists of 19 leaves which are terminal nodes and our tree starts with the most important features of the data set, the policyholder age. Furthermore the root consists of 4200 samples obtained based on the ratio between a training and validation sets (i.e. 70:30). As one can see the model is easy to interpret and enjoys high accuracy.

Acknowledgment

Authors are grateful for supports received from Iran Insurance Research Center (IRC).

References

- [1] Compulsory Third Party Liability Insurance Damage due to Vehicle Accidents (2015) .
- [2] A.C. Yeo, K.A. Smith, R.J Willis, M. Brooks. Clustering technique for risk classification and prediction of claim costs in the automobile insurance industry. *Intelligent Systems in Accounting, Finance & Management*. 2001 Mar;10(1):39-50.
- [3] C.K Fan, W.Y. Wang. A comparison of underwriting decision making between telematics-enabled UBI and traditional auto insurance. *Advances in Management and Applied Economics*. 2017;7(1):17.
- [4] J. Son, M. Park, B.B. Park. The effect of age, gender and roadway environment on the acceptance and effectiveness of Advanced Driver Assistance Systems. *Transportation research part F: traffic psychology and behaviour*. 2015 May 1;31:12-24.
- [5] Y. Bian, C. Yang, J.L. Zhao, L. Liang. Good drivers pay less: A study of usage-based vehicle insurance models. *Transportation research part A: policy and practice*. 2018 Jan 1;107:20-34.
- [6] W.Y. Loh. *Classification and regression trees*. Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery. 2011 Jan;1(1):14-23.

- [7] C. Machuca, M.V. Vettore, M. Krasuska, S.R. Baker, P.G. Robinson. Using classification and regression tree modelling to investigate response shift patterns in dentine hypersensitivity. *BMC medical research methodology*. 2017 Dec;17(1):120.
- [8] Y.Y. Song, L.U. Ying. Decision tree methods: applications for classification and prediction. *Shanghai archives of psychiatry*. 2015 Apr 25;27(2):130.
- [9] V. Podgorelec, P. Kokol, B. Stiglic, I. Rozman. Decision trees: an overview and their use in medicine. *Journal of medical systems*. 2002 Oct 1;26(5):445-63.
- [10] L. E. Brandão, J.S. Dyer, W.J. Hahn. Using binomial decision trees to solve real-option valuation problems. *Decision Analysis*. 2005 Jun;2(2):69-88.
- [11] M. Zięba, S.K. Tomczak, J.M. Tomczak. Ensemble boosted trees with synthetic features generation in application to bankruptcy prediction. *Expert Systems with Applications*. 2016 Oct 1;58:93-101.
- [12] C.S Huang, Y.J. Lin, C.C Lin. Implementation of classifiers for choosing insurance policy using decision trees: a case study. *WSEAS Transactions on Computers*. 2008 Oct 1;7(10):1679-89.
- [13] M.H. Spekkers, M. Kok, F.H. Clemens, J.A. Ten Veldhuis. Decision-tree analysis of factors influencing rainfall-related building structure and content damage. *Natural Hazards and Earth System Sciences*. 2014 Sep 24;14(9):2531-47.
- [14] Z. Quan , E.A. Valdez. Predictive analytics of insurance claims using multivariate decision trees. *Dependence Modeling*. 2018 Dec 1;6(1):377-407.
- [15] C.C. Chern, Y.J. Chen, B. Hsiao. Decision tree–based classifier in providing telehealth service. *BMC medical informatics and decision making*. 2019 Dec;19(1):104.
- [16] S.S. Thakur and J.K. Sing, Prediction of Online Vehicle Insurance System using Decision Tree Classifier and Bayes Classifier–A Comparative Analysis. *International Journal of Computer Applications*, 975(2014), p.8887.
- [17] Thakur SS, Sing JK. Mining Customer's Data for Vehicle Insurance Prediction System using Decision Tree Classifier. *International Journal on Recent Trends in Engineering & Technology*. 2013 Jul 1;9(1):121
- [18] N.K Frempong, N. Nicholas, M.A. Boateng. Decision tree as a predictive modeling tool for auto insurance claims. *Int. J. Statist. Appl.*. 2017;7(2):117-20.

Email: khanizadeh@irc.ac.ir

Email: oskoei@atu.ac.ir

Deep Neural Networks with Long Short-Term Memory for Human Mortality Modeling

Jose Garrido¹

Department of Mathematics and Statistics, Concordia University, Montreal, Canada

Abstract

Accurate modeling and forecasting of human mortality rates is important in actuarial science, to price life insurance products, pension plan evaluations, and in finance, to price derivative products used to hedge longevity risk. Data shows that mortality rates have been decreasing at all ages over time, especially in the last century. Predicting the extent of future longevity improvement represents a difficult and important problem for the life insurance industry and for sponsors of pension plans and social security programs.

The most popular methodology to forecast future mortality improvement was proposed by Lee and Carter (1992, JASA). It consists of a two-steps process, shown to suffer from identifiability issues, both in the Lee-Carter Model and its subsequent extensions, mostly due to the inherent two-steps model setup. We propose a very distinct, data-driven approach using a class of Deep Neural Networks to model and forecast human mortality. The main component in the neural networks is a long short-term memory (NNLS-TM) layer, which was introduced by Hochreiter and Schmidhuber (1997, NC), to fix vanishing gradients in simple recurrent neural networks. The model can be constructed for short-term as well as for long-term forecasting, respectively.

We model the dependence mortality improvement observed simultaneously in different countries. Current mortality improvement models are fitted to single country sub-populations separately, even if improvement trends are similar in different countries. The multi-population problem presents serious computational challenges that we tackle with NNLS-TMs, fitted to learn from single country populations included in the Human Mortality Database (<https://www.mortality.org/>).

(This is joint work with Ran Xu, Department of Mathematical Sciences, Xi'an Jiaotong-Liverpool University, Suzhou, China)

Email: jose.garrido@concordia.ca

¹speaker

Risk Identification with ERM approach in an Insurance Company: Taavon Insurance Company Case Study

Maryam Salmasi¹
Taavon Insurance Company, Tehran, Iran

Abstract

Any choices we make in the pursuit of objectives has its risks and all companies must manage their risks to reach long term success. In this short and descriptive industrial applied article, we review the suitable risk management approach for an Iranian insurance company and its experiences in identifying different risks at different levels which are strategic, business and process and review some examples. We also follow the definition and different frameworks to implement an Enterprise Risk Management (ERM), the importance of it, and a short review of the regulation about risk management in the Iranian insurance industry.

Keywords: Risk Identification, Enterprise Risk Management, Insurance

Introduction

Any choice we make in the pursuit of objectives has its risks. From day to day operational decisions to the fundamental trade-offs, dealing with risk is a part of decision-making. As COSO says, good risk management and internal control are necessary for long term success of all organizations [3]. In this article we review the importance, definition and different approach to implement an Enterprise Risk Management (ERM) with focus on Taavon Insurance Company experience on implementing a customized ERM approach.

Taavon insurance co. has established since 2013 with the primary aim of joining the cooperative sector in the field of banking and insurance and also access to monetary and financial markets of the country. After obtaining the principal agreement of the Central Insurance of Iran and the Ministry of Cooperatives, has become the first cooperative-public-corporation company. Today, after implementing numbers of improvement projects related to customer satisfaction such as ISO9001 in quality management, ISO10002 in complaint handling and ISO10004 in customer satisfaction assessment, the company is following a growing trend from 2019. This company want to be excellence at providing distinct and innovative insurance services for its customer. Moreover, the vision is to be the first choice of customers, sales network and professionals. With this introduction, the article will review the risk management role in this company.

1.1 Definitions

Risk is the effect of uncertainty on objectives which can have different aspects and categories, and can be applied at different levels [1].

Risk management is to coordinate activities to direct and control an organization with regard to risk [1].

¹ PhD Student of Management at Kharazmi University; Chef Risk Officer at Taavon Insurance Company.

Risk source is an element which alone or in combination has the potential to give rise to risk [1].

Event is an occurrence or change of a particular set of circumstances. An event can have one or more occurrences, and can have several causes and several consequences. Furthermore, it is notable that an event can be a risk source [1].

Consequence is the outcome of an event affecting objectives [1].

Risk is usually expressed in terms of risk sources, potential events, their consequences and their likelihood [1].

Enterprise Risk Management (ERM) is the process or discipline by which organizations in all industries assess, control, exploit, finance, and monitor risks from all sources for the purpose of increasing the organization's short and long term value to its stakeholders [2].

1.2 The Importance of Risk Management

The purpose of risk management is the creation and protection of value. It improves performance, encourages innovation and supports the achievement of objectives [1]. It is reasonable to expect that the forces like more complicated risks, external pressure and changes never stop in a company environment. Accordingly, risk management practices will become more and more sophisticated. As capabilities continue to improve, organizations will increasingly adopt ERM [2]. Risks can be found on different levels, actions and decisions of a company including an insurance company.

Several texts have discussed concepts such as "strategic risk management", "integrated risk management" and "holistic risk management". These concepts are similar to, even synonymous with, ERM in that they all emphasize a comprehensive view of risk and risk management, and the view that risk management can be a value-creating, in addition to a risk-mitigating, process [2].

An insurance company is inherently established to manage the risk of its customers but the question is whether it is aware of its own risks and manages them well or no; and how an insurance company as a firm should do that well enough.

As it is mentioned Taavon Insurance Co. as an insurance company inherently is expert in managing the risk of its customers. This background provides a good platform for implementing an ERM.

But in order to implement a new system, all companies need to choose the appropriate approach, plan and train well enough, develop the culture and formulate relevant procedures within themselves. So here's a look at what Taavon has done so far.

The need to attention to corporate governance and risk management in Iranian insurance companies is new. Central Insurance of Iran as the highest regulatory body of insurance in Iran has just announced requirements for corporate governance including risk management to insurance companies in 2017. Central Insurance noticed rules include necessity to establish risk committee under the supervision of the board of directors and risk management department in all insurance companies as it mentioned in regulations No.90. Insurance companies should also assess their capital adequacy within 4 different risks of: 1. underwriting risk (R1), 2. market risk (R2), 3. credit risk (R3), and 4. liquidity risk (R4) as it mentioned in solvency regulation No 69. The formula of solvency rate is as follows (1) and (2), which is better to be upper than 100%.

$$\text{Required Capital (RC)} = \sqrt{(R_1^2 + R_2^2 + R_3^2 + R_4^2)}, \quad (1)$$

$$\text{Solvency Rate} = \frac{\text{Current Capital}}{\text{RC}} \times 100. \quad (2)$$

In last assessment Taavon has reached 138%. This rate puts this company in the first level of solvency, which means has a well control on its financial risks.

As you may concern, calculating the solvency and capital adequacy is not enough to establish an ERM in an insurance company. So the main question is what kind of model is appropriate for Taavon Insurance co. considering the internal and external affecting activities factors.

1.3 Enterprise Risk Management Frameworks

There are numbers of ERM frameworks such as CAS, COSO, ISO31000, Standard and Poor's (S&P), New York Stock Exchange (NYSE), and Sarbanes-Oxley Act and etc, each of which describes an approach for identifying, analyzing, and responding to risks which are facing the enterprise and its objectives. Some important frameworks which Taavon Insurance Co. has considered are: 1) Casualty Actuarial Society framework (CAS); 2) Committee of Sponsoring Organizations of the Treadway Commission ERM framework (COSO ERM 2017), it is notable that COSO which is published in 2004 has major differences with the new version which is published in 2017; and 3) ISO31000: 2018.

With a benchmark on "Coldhard Steel" ERM Framework, Casualty Actuarial Society Enterprise Risk Management Committee has offered a four categories framework for risk management on 2003. With reference to that framework, in general, enterprises are exposed to risks that can be categorized into the following four types. The precise slotting of individual risk factors under each of these four categories is less important than the recognition that ERM covers all categories and all material risk factors that can influence the organization's value [2]. The framework also introduces a process to manage risks.

- Hazard Risks include risks from:
 - fire and other property damage,
 - windstorm and other natural perils,
 - theft and other crime, personal injury,
 - business interruption,
 - disease and disability (including work-related injuries and diseases), and
 - liability claims.
- Financial Risks include risks from:
 - price (e.g. asset value, interest rate, foreign exchange, commodity),
 - liquidity (e.g. cash flow, call risk, opportunity cost),
 - credit (e.g. default, downgrade),
 - inflation/purchasing power, and
 - hedging/ basis risk.
- Operational Risks include risks from:
 - business operations (e.g., human resources, product development, capacity, efficiency, product/service failure, channel management, supply chain management, business cyclicity),
 - empowerment (e.g., leadership, change readiness),
 - information technology (e.g., relevance, availability), and
 - information/ business reporting (e.g., budgeting and planning, accounting information, pension fund, investment evaluation, taxation).
- Strategic Risks include risks from:
 - reputational damage (e.g., trademark/brand erosion, fraud, unfavorable publicity)
 - competition,
 - customer wants,
 - demographic and social/ cultural trends,
 - technological innovation,
 - capital availability, and
 - regulatory and political trends.

New COSO Framework focused on five components and 20 key principles within each of the five components. The five components are:

- Governance and Culture,
- Strategy and Objective Settings,
- Performance,
- Review and Revision,
- Information, Communication and Reporting.

The previous version of COSO which was introduced in 2004 had four objectives categories which are comparable to CAS framework:

- Strategy - high-level goals, aligned with and supporting the organization's mission,
- Operations - effective and efficient use of resources,
- Financial Reporting - reliability of operational and financial reporting,
- Compliance - compliance with applicable laws and regulations.

ISO 31000:2018, Risk management – Guidelines, provides principles, framework and a process for managing risk. It can be used by any organization regardless of its size, activity or sector. This guideline does not provide a clear classification of risks but has provided a simple and understandable framework for risk management.

1.4 Taavon Insurance Experience

Taavon Insurance co. has started the implementation of a customized ERM approach in 2019. Some of the important decision factors for Taavon to select and design its own ERM model are comprehensiveness, precision, easiness of implementation and providing clear risk categories. Because of maturity enterprise risk management level and being new, the last two factors had weighted more.

Because of the following reasons in Table 1 Taavon took the CAS approach to classify the business level risks and also chose the ISO31000 approach to implement the framework.

Approach	Comprehensiveness	Precision	Easiness	Risk Categories
CAS	***	*	**	+
COSO 2004	**	*	*	+
COSO 2017	***	*	*	-
ISO 31000	*	*	***	-

Table 1: Taavon Framework Choosing Factors

The insurance company should define the scope of its risk management activities. As the risk management process may be applied at different levels (e.g. strategic, operational, program, project, or other activities), it is important to be clear about the scope under consideration, the relevant objectives to be considered and their alignment with organizational objectives [1].

The main goals of Taavon risk management system is to increase solvency level, improving claims ratio and increasing customer satisfaction beside continuation of the business activity protecting shareholder rights. By modeling all the frameworks, Taavon decided to develop a three level risk management framework. It has chosen to define the scope of its risk management in three levels of strategic, business, and process.

- Strategic Level Risks: Occurrence of these risks, make huge and fundamental changes in the activities of the company such as suspension of activity or losing a significant share of profits or even dissolution. At least once a year, they are identified and dealt with by the Strategic Planning Committee
- Business Level Risks: If these risks occur, interfere with the execution of operations of a department or departments or set of systems. For developing these level of risk analysis, Taavon uses CAS categories framework. Business level risks are identified at least once a year by interviewing risk owners, then are analyzed, assessed and decided on their exposure by risk committee.
- Process Level Risks: If these kind of risks occur, they may interfere with the execution of one process. They are identified, managed, and updated by the owner of each process.

With this approach Taavon has planned and fulfilled some goals such as: formed an organizational department for risk management; formed the risk committee; formulated and approved the policy statement for risk management; formulated and approved the risk management process and procedure; held interviews with risk owners to identify business level risks; identified all process level risk; determined the risk representative for each department; reporting different risk factors; and many other actions which are all of which have been instrumental in developing a risk management culture and day to day managing actions.

As it mentioned before Taavon process to manage risks is based on ISO31000 which makes it a seven steps process as follows in figure 1.

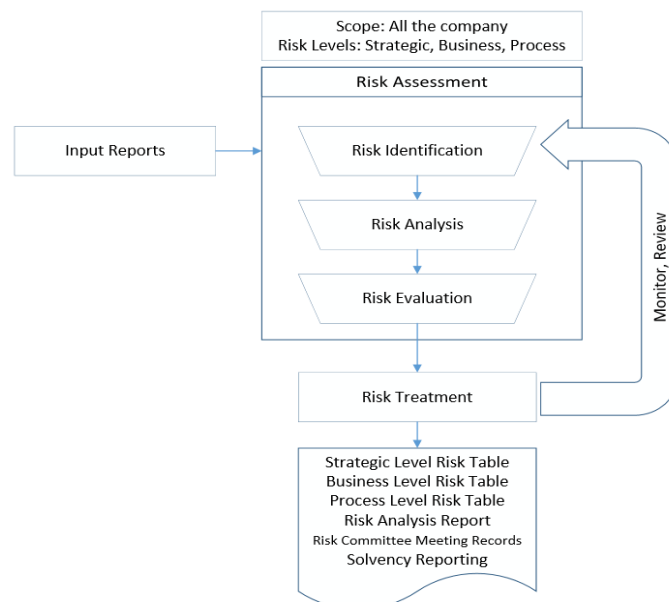


Figure 1: Taavon Risk Management Process

Outcomes

In this article the results are focused on the second step of the process which is “Risk Identification” in three levels of strategic, business and process.

The purpose of risk identification is to find, recognize and describe risks that might prevent an organization achieving its objectives. The organization should identify risks, whether or not their sources are under its control. Consideration should be given that there may be more than one type of outcome, which may result in a variety of tangible or intangible consequences [1]. In this step relevant, appropriate and up-to-date information is important in identifying risks. The organization can use a range of techniques for identifying uncertainties that may affect one or more objectives.

As it is defined in this company different committee or different authorities are responsible to identify risk due to its level. Strategic level risks are identified by the Strategic Planning Committee once a year; business level risks are defined by Chief Risk Officer (CRO) with reference to the interview results with risk owners and are confirmed by Risk Committee and it may be categorized as hazard risk, financial risk, operational risk, strategic decision making risks; and process level risks are identified by the process owner and confirmed by CRO.

Furthermore, it is very important to note that to identify any risk or its changes, good and suitable sources of information are necessary. Here are three different examples of Taavon risk identifications and its used data sources in table 2 to 4. Also note that as it was mentioned in risk definition, we usually expressed “risk” in terms of 1) risk sources, 2) potential events, and 3) their consequences; which are clarified at identification step.

Information	Risk Level	Strategic
	Input Information	Central Insurance negative score report interview with CEO and board of directors media reports others companies case study
	Related Objective	continuation of the business activity
Risk	Risk sources	failure to comply with regulations not paying attention to the Central Insurance hints
	Potential events/ Risk	suspension of activity
	Consequences	brand infamy losing customers bankruptcy

Table 2: Taavon Identified Strategic Risk Sample

Information	Risk Level	Business
	CAS category	Operational Risks
	Input Information	staff satisfaction annual report interview with all managers previous experiences results of key performance indicators such as comparison of average salaries
	Related Objective	sales increase appropriate staff loyalty
Risk	Risk sources	unsatisfied staffs
	Potential events/ Risk	loosing key persons
	Consequences	loosing knowledge increase staffing costs losing customers and sales decline

Table 3: Taavon Identified Business Risk Sample

Information	Risk Level	Process
	Process Name	Health Insurance Process
	Input Information	customer satisfaction report process analysis results of process indicators such as duration of claim processing internal quality audit report
	Related Objective	speed of service delivery
Risk	Risk sources	lack of control on service delivery duration incomplete data recording in the software system lack of cooperation from loss assessor process defects unsatisfied staffs
	Potential events/ Risk	increasing duration of claim processing
	Consequences	unsatisfied customer losing customers and sales decline

Table 2: Taavon Identified Process Risk Sample

However, the concept of risk management has been introduced to the world based on the principles of insurance and banking, in the Iranian insurance industry, enterprise risk management is a new concept. Therefore, beside Iranian insurance companies must work to improve and deepen their laws and regulations, they should try implementing these concepts through innovation and by applying latest models to their organizations. In this brief review, we review the customized model of ERM and the critical phase of risk identification in Taavon Insurance Co.

References

- [1] ISO/TC 262 Risk management Technical Committee, *ISO 31000:2018 Risk management — Guidelines*, Feb 2018, Edition: 2.
- [2] CAS: Casualty Actuarial Society: Enterprise Risk Management Committee, *Overview of Enterprise Risk Management*, May 2003, <https://www.casact.org/area/erm>.
- [3] COSO: Committee of Sponsoring Organizations of the Treadway Commission, *Enterprise Risk Management, Integrating with Strategy and Performance*, June 2017.

Email: m-salmasi@taavon-ins.ir, std_salmasi@khu.ac.ir

Insurance Pricing: from theory to reality

Rahim Mahmoudvand¹
Bu-Ali Sina University, Hamedan, Iran

Abstract

Risk is an inevitable part of the world that can be found in different forms around us. We may be able to control the negative effect of some risks. But, insurance is an appropriate method that can be used for treating many different types of risks.

Insurance pricing is a pivot in insurance industry as it is associated to all stakeholders. Insured, insurers and other stakeholders of the insurance industry follow up insurance pricing carefully. However, in most cases, the insurers implement insurance pricing. The price of insurance is normally a function of the cost of production. Unlike of many products, the costs of insurance products are not fixed and depend on a range of factors.

In this paper, I have investigated the mechanisms by which these factors associate with insurance pricing. Then, I tried to identify drawbacks of these mechanisms.

Keywords: Insurance, Risk, Modeling.

Mathematics Subject Classification [2018]: 91B30, 62P05

1 Introduction

We all understand the meaning of risk but surprisingly there is no definition for risk that is agreed by all researchers. A relatively common definition for risk says that it is a condition in which there is a possibility of an adverse deviation from a desired outcome that is expected or hoped for; see Voughan and Voughan (2013).

Core feature of insurance is sharing risk in exchange for payment. Insurance pricing is a methodology for determining the price of risk for insured. An appropriate pricing ensure that insurance company set fair and adequate premium given the competitive nature. Simplicity and stability are two other criteria that assumed to be satisfied in insurance pricing. In my opinion, an important question in insurance pricing is about the meaning of “fairness”. Our understanding by fairness is driven by both culture and legislation. From viewpoint of insurers, the fairness tied up with cost of the insurance product. In this paper, I want to discuss about the cost of the insurance product

2 Main results

Let me start with the list of costs in insurance pricing. There might be a long list for the cost of insurance, but these are categorized as below:

- Losses and loss adjustment expenses,
- Acquisition expenses,
- Administrative expenses,

¹ speaker

- Taxes,
- Profits and contingencies.

Actuaries commonly focus on the first item. This can be understood from main actuarial journals. For instance, I checked four famous actuarial journals (see Table 1). First, I found the total number of publications from Scopus (<https://www.scopus.com/sources?zone=TopNavBar&origin=sourceinfo>) for each journal. Then, I decided to see all publications in each journal. Unfortunately, I couldn't access to all contents so I just tried to look at their recent publications. My conclusion is that the main focus of the publications is on losses and loss adjustment expenses. One can complete this checking to find an exact estimate of the percentage of the publications on losses and loss adjustment expenses.

Table 1: A sample of actuarial journals

Journal	Starting year	Number of papers
Scandinavian Actuarial Journal	1918	1964
ASTIN Bulletin	1958	1247
Insurance Mathematics and Economics	1992	2051
North Actuarial American Journal	1997	959

The other sources of theories are books. I also checked several books in the area of actuarial science. Figure 1 show four famous actuarial books. Again the conclusion was similar to the results that I got by review of the journals.

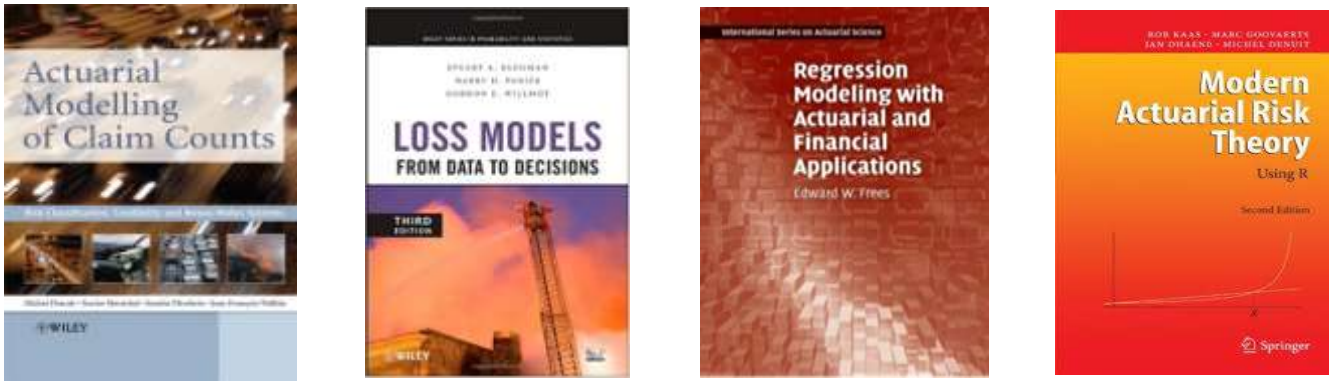


Figure 1: A sample actuarial book

Let us to see what theories we use for insurance pricing. Denote by random variable X the total incurred losses with an insurable risk. Then, π_X , that is called “premium”, will be defined as below:

$$\pi_X : S_X \rightarrow \mathbb{R}^+, \quad (2)$$

where S_X is the support of r.v. X . Finding π_X is based on the assumption that a contingent claim expenses can be compensated by fixed payments which is indeed the premium. Several common forms for π_X are given in Table 2.

Table 2: Common forms of premium principles

Name	Formula	condition
Expected value principle	$(1 + \rho) \times E(X)$	$\rho > 0$
Standard deviation principle	$E(X) + \alpha \times \sigma(X)$	$\alpha > 0$
Variance principle	$E(X) + \beta \times \sigma^2(X)$	$\beta > 0$
Zero utility principle	$\pi_X = \pi; \text{ where } E[u(\pi - X)] = u(0)$	

In these approaches, mean and variance of X cover incurred losses and parameters ρ , α and β applied to capture other costs of the insurance product (items 2-5 in above mentioned list). In order to find mean, variance and higher

moments of X , we require knowing loss distribution. There is a vast literature on this topic. The above mentioned books include many good references. Studies on determining π_X can be divided into two categories:

- Structural form of π_X and their characteristics,
- Computational issues.

Theoretically, a premium principle is desirable for insurer if it develops an adequate premium return. However, literature shows that most of business lines prefer to use expected value principle. Almost all life insurance products use expected value principle. Posterior rate making such as Bonus-Malus Systems are indeed based on the expected value principle. I think the main reason for this choice is the challenges that we have in practice when we want to use other principles. By the way, simplicity and flexibility of expected value principle provide several advantages. For instance, using GLM and Bayes theorem one can improve the level of fairness with this principle. These tools can be used to model variability of losses by risk factors. Such aims can be indeed achieved when we have access to data. Nevertheless, the problem of the creditability of insurance data is one of the most fundamental problems for the non-life actuary. But in the last few years, the impact of big data on the assessment of individual risks results in a growing debate on what is a fair actuarial price; see. e.g. Martinez et al, 2016.

Now, let me give a brief discussion about other sources of insurance cost. Acquisition expenses are agent's commission which is very common in life insurance. This cost is not stochastic, but it could be reduced by new technologies partly. Administrative expenses include costs other than losses and acquisitions and seem to be fixed. But in practice these costs depends on the quality of insurer's workforce. Taxes are percentage of the total premium, so higher costs produce higher taxes. The last item is a hoped-for return on capital.

Let come back again to the expected value principle. The parameter ρ included in this premium principle to capture all costs other than losses and adjustment losses. The above discussion on the insurance costs show that ρ vary by company for the same insurable risk. So, the total premium for a unique risk may differ dramatically and this is in contradiction with fairness.

References

- [1] Antonio José Heras Martínez, David Teira, Pierre-Charles Pradier. *What was fair in acturial fairness?*. 2016. halshs-01400213.
- [2] H. Buhlmann, *Mathematical Methods in Risk Theory*, 2nd ed., Springer, Berlin, 1996.
- [3] E. J. Vaughan, T. M. Vaughan, *Fundamental of Risk and Insurance*, 11th ed., John Wiley, 2013.

Email: r.mahmodvand@gmail.com

Solving Parametric Fractional Differential Equations Arising from Rough Heston Model using Quasi-Linearization and Spectral Collocation

Ali Foroush Bastani¹

Department of Mathematics, Institute for Advanced Studies in Basic Sciences, Zanjan, Iran

Abstract

The rough Heston model has recently attracted the attention of many finance practitioners and researchers as it maintains the basic structure of the classical Heston model while having descriptive capabilities in terms of micro-structural foundations of the market. Using the fact that the characteristic function of log-price in this model could be expressed in terms of the solution of a nonlinear parametric fractional Riccati differential equation not admitting a closed-form solution, devising efficient numerical schemes for pricing and calibration under this model has become a crucial need in the computational finance community. Although the fractional Adams method has been used in most of the recent studies on the rough Heston model, this method suffers from some stability and convergence issues in treating the problem. In this paper, we present a numerical method based on Newton-Kantorovich quasi-linearization to solve the nonlinearity issue followed by spectral collocation based on “poly-fractionomials” to approximate the fractional derivatives in an accurate and efficient manner. We provide sufficient conditions under which our method is convergent and the order of convergence is also obtained. In order to guarantee the specified convergence rate, we first prove some regularity results on the linearized problem and then employ the proposed scheme to solve a practical calibration problem from the SPX options market. The efficiency of the proposed method is illustrated by comparing the results with the fractional Adams method.

Keywords: Rough Heston model, Fractional nonlinear Riccati differential equation, Spectral collocation

Email: bastani@iasbs.ac.ir

¹speaker

An efficient pricing method for basket options under jump diffusion model

Ali Safdari-Vaighani¹

Department of Mathematics, Allameh Tabataba'i University, Tehran, Iran.

1 Introduction

Option contracts under actual market conditions which are more complex than a simple Black-Schole model are important hedging strategies in the modern financial market. Basket options have an important role in the FX market in particular as they offer protection against drops in all the currencies at the same time. Basket options provide a cheaper alternative to buying individual options on each asset to hedge against risk and also the transaction costs are greatly lowered when one only buys a single option rather than multiple options. Basket options are attractive products which required the reliable pricing method to take all the beneficial characteristics of a basket option such as correlation effect of underlying assets.

The leptokurtic feature has been observed since 1950's. However classical finance models simply ignore this feature. For example, in the Black-Scholes Brownian motion model, the stock price is modeling as

$$dS(t) = \mu S(t)dt + \sigma S(t)dW_t$$

In this model, the continuous compounded return, $\ln(S(t)/S(0))$, has a normal distribution, which it is not consistent with "leptokurtic feature".

After the 1987 market crash, the implied volatility smile becomes economically significant and the performance of the Black-Scholes model deteriorated [Rubinstein (1985)]. Many studies have been conducted to modify the Black-Scholes model to explain the three empirical stylized facts, namely the leptokurtic feature, volatility clustering effect, and implied volatility smile.

- Fractal Brownian motions models [Rogers (1997)]
- Models based on Levy processes [Cont and Tankov (2004)]
- Stochastic volatility models [Fouque et al. (2000), Heston (1993)]
- Jump-diffusion models [Merton (1976), Kou (2002)]

Assume that the asset price $S_i, i = 1 \dots d$ follows the process

$$\frac{dS_i}{S_i} = \mu dt + \sigma_i dW_i + (e^{J_i} - 1)dq, \quad (1)$$

¹speaker

Using the Ito's formula, the contingent claim $V(S, t)$ that depends on $S = (S_1, \dots, S_d) \in \tilde{\Omega} = \mathbb{R}_+^d$ can be derived by taking the expectation under the risk-neutral process.

$$\begin{aligned} \frac{\partial V}{\partial t}(S, t) = & -\frac{1}{2} \sum_{i=1}^d \sum_{j=1}^d \rho_{ij} \sigma_i \sigma_j S_i S_j \frac{\partial^2 V}{\partial S_i \partial S_j}(S, t) - \sum_{i=1}^d r S_i \frac{\partial V}{\partial S_i}(S, t) + rV(S, t) \\ & - \lambda \int_{\tilde{\Omega}} [V(Se^J, t) - V(S, t) - \sum_{i=1}^d S_i (e^{J_i} - 1) \frac{\partial V}{\partial S_i}(S, t)] g(J) dJ \end{aligned} \quad (2)$$

Let $S_i = e^{x_i}$ and $\tau = T - t$ and by the variable transformations, $V(e^x, T - \tau) = U(x, \tau)$ the equation can be rewrite in the more tractable form:

$$\begin{aligned} \frac{\partial U}{\partial \tau} = & \underbrace{\frac{1}{2} \sum_{i=1}^d \sum_{j=1}^d \rho_{ij} \sigma_i \sigma_j \frac{\partial^2 U}{\partial x_i \partial x_j} + \sum_{i=1}^d (r - \frac{\sigma_i^2}{2} - \kappa_i \lambda) \frac{\partial U}{\partial x_i}}_{\mathcal{LU}} - (r + \lambda)U \\ & + \lambda \underbrace{\int_{\Omega} U(x + J, \tau) g(J) dJ}_{\mathcal{IU}} \quad (x, \tau) \in \Omega \times (0, T], \end{aligned} \quad (3)$$

where $\kappa_i = \int_{\Omega} (e^{J_i} - 1) g(J) dJ$ with the distribution function of jumps $g(J)$ and λ is mean arrival rate of jumps. The resulting PIDE poses difficulties to solve numerically due to

- We are dealing with d-dimensional PIDE problem.
- Truncating the infinite domain of the PIDE to finite boundaries.
- The payoff function possesses a discontinuity in its first derivative at the exercise price.
- We should avoid producing a dense system of equations when we discretize the equation.

Basket options evaluation are not tied to Monte-Carlo simulation method. They have indeed been actively researched by other numerical methods. We list the following four popular methods:

- Finite difference scheme for the two dimensional PDE (Nielsen et. al. 2000), the two dimensional PIDE (Forsyth et.al. 2008)
- Radial basis function approximation method for the two dimensional PDE (Pettersson et.al. 2008, Safdari et.al. 2015)
- Fast Fourier transform (FFT) approach, (Oosterlee et.al. 2008)
- Monte-Carlo simulation (Glasserman 2004)

2 Radial basis function partition of unity

In a partition of unity (PU) scheme, local approximations on overlapping patches that form a cover of the computational domain are weighted together by compactly supported partition of unity weight functions to form the global approximation. Let $\{\Omega_j\}_{j=1}^M$ are overlapping patches that form a cover of the computational domain Ω . The global approximation function $s(x)$ in domain Ω to the solution function $u(x)$ is constructed as

$$s(x) = \sum_{j=1}^M w_j(x) s_j(x), \quad (4)$$

where $s_j(x)$, $j = 1, \dots, M$ are local interpolates and $s_j(x_i) = u(x_i)$ for each node point $x_i \in \Omega_j$

$$s_j(x) = \sum_{k \in J(\Omega_j)} \psi_k(x) u_k,$$

The partition of unity weight functions w_j , are constructed using Shepard's method

$$w_j(x) = \frac{\varphi_j(x)}{\sum_{k \in I(x)} \varphi_k(x)},$$

which $\varphi_j(x)$ are compactly supported functions with support on Ω_j . The global approximation function for time depended problem can be introduced as

$$s(x, t) = \sum_{j=1}^M w_j(x) s_j(x, t), \quad (5)$$

Why we have chose the RBF-PUM as numerical method

- Easy to implement in any number of dimensions
- Allows for local adaptivity. Patches can be locally refined and have shapes adapted to the local solution behavior.
- Produce the sparse differentiation matrices
- Leads to the high order of convergence rate

References

- [1] S. S. CLIFT AND P. A. FORSYTH, *Numerical solution of two asset jump diffusion models for option valuation*, Appl. Numer. Math. **58**:6 (2008), 743–782.
- [2] R. CONT, P. TANKOV, *Financial Modelling with Jump Processe*, CRC Financial Mathematics Series, Chapman and Hall, 2004.
- [3] R. C. MERTON, *Option pricing when underlying stock returns are discontinuous*, Journal of financial economics **3**:1-2 (1976), 125–144.
- [4] U. PETTERSSON, E. LARSSON, G. MARCUSSEON AND J. PERSSON, *Improved radial basis function methods for multi-dimensional option pricing*, J. Comput. Appl. Math. **222**:1 (2008), 82–93.
- [5] A. Safdari-Vaighani, A. Heryudono and E. Larsson, *A radial basis function partition of unity collocation method for convection-diffusion equations arising in financial applications*, J. Sci. Comput. **64**:2 (2015), 341–367.

A jump diffusion model when both underlying and volatility contain correlated jumps

F. Soleymani¹

Department of Mathematics, Institute for Advanced Studies in Basic Sciences (IASBS), Zanjan 45137-66731, Iran

Abstract

An option pricing model when not only the underlying asset but also the volatility are coming from stochastic drivers along with contemporaneous jumps is considered in this work. To price under this introduced 1+2 dimensional PIDE problem, a local scheme will be constructed to be fast and efficient by relying heavily on sparse matrices and adaptive node layouts. Several issues are discussed to circumvents on the challenges of pricing under such a jump-diffusion model.

Keywords: Stochastic volatility model; double integral; jump diffusion; contemporaneous jumps
Mathematics Subject Classification [2018]: 65M22; 91B25

1 Introduction

The standard option theory in finance assumes that the logarithm of asset price is normally distributed, [3]. However, in practice, the observed distributions are not normal - they exhibit ‘fatter tails’, i.e. the probability of very large moves in either direction is larger than allowed by normal distribution. Therefore, the jump-diffusion method has been brought to provide plausible mechanism for explaining why fat tails exists and what are their consequences, [2].

Both cases, when jumps in volatility and returns are occurring independently (known as SVIJ), as well as when they are taking place contemporaneously (known as SVCJ), have been discussed in the work [3]. To illustrate further, both SVIJ and SVCJ are good models for option pricing under jump diffusions, but the SVCJ model could provide better fit to the observations of the market, see also [4].

There are several works handling the pricing problems which need a PIDE’s solution in two spatial dimensions (specially with a double integral source). The work [1] investigated a finite element (FE) approach for solving this problem. In [7], an implicit–explicit computational scheme is investigated for solving the Bates SV model [2] using an operator splitting.

In this work, we consider a new convex combination of the two already well-known RBFs, viz, the multiquadric (MQ) and the inverse multiquadric (IMQ) RBFs as follows:

$$\phi(r_i) = \theta(c^2 + r_i^2)^{\frac{1}{2}} + (1 - \theta)(c^2 + r_i^2)^{-\frac{1}{2}}, \quad \theta \in [0, 1], \quad i = 1, 2, \dots, m, \quad (1)$$

where $r_i = \|\mathbf{x} - \mathbf{x}_i\|_2$ denotes the Euclidean distance and the parameter of shape is c . This RBF is denoted by CRBF as the convex combination of the aforementioned RBFs.

Here the motivation is to contribute in solving a (1+2)D PIDE originated by the model of the SVCJ consisting of a multiple integral. We compute and use the weights of the novel CRBF–FD scheme for *spatially* discretizing the PIDE problem. A Krylov subspace method alongside nonuniform discretizations provide an efficient tool for tackling the SVCJ model in computational finance.

A main idea of this work is when we compute the double integral resulting from the jump terms, the solution values in each discretized cell is approximated by the average of the values in the four corners. This contribution allows for analytic integration of the jump density which is novel in literature.

¹speaker

2 The SVCJ PIDE

Under the dynamic of the SVCJ model, for pricing in case of a European option, we have the following PIDE which is formulated in forward time [5]:

$$\begin{aligned} \frac{\partial u(s, v, \tau)}{\partial \tau} &= \frac{1}{2} v s^2 \frac{\partial^2 u(s, v, \tau)}{\partial s^2} + \frac{1}{2} \sigma^2 v \frac{\partial^2 u(s, v, \tau)}{\partial v^2} + \rho \sigma v s \frac{\partial^2 u(s, v, \tau)}{\partial s \partial v} \\ &\quad + (r - q - \lambda \xi) s \frac{\partial u(s, v, \tau)}{\partial s} + \kappa(\varpi - v) \frac{\partial u(s, v, \tau)}{\partial v} - (r + \lambda) u(s, v, \tau) \\ &\quad + \lambda \underbrace{\int_0^\infty \int_0^\infty u(s z^s, v + z^v, \tau) \mathbf{p}(z^s, z^v) dz^v dz^s}_{\mathcal{L}_I u(s, v, \tau)} \\ &=: \mathcal{L}u(s, v, \tau), \quad s, v \geq 0, \end{aligned} \quad (2)$$

wherein $\tau = T - t$ and $\mathcal{L}u(s, v, \tau) = \mathcal{L}_D u(s, v, \tau) + \lambda \mathcal{L}_I u(s, v, \tau)$. The 2D probability density function (PDF) $\mathbf{p}$ has a log-normal distribution $\mathbf{p}(z^s, z^v)$ for z^s and normal distribution along z^v , is expressed as [3]:

$$\mathbf{p}(z^s, z^v) = \frac{1}{\sqrt{2\pi} z^s \delta \nu} e^{\left(-\frac{z^v}{\nu} - \frac{(\ln(z^s) - \gamma - \rho_J z^v)^2}{2\delta^2} \right)}. \quad (3)$$

3 Nodes layout

For producing a mesh to be as coarse as possible away from the hot area and as refined as possible in the hot zone, some strategies have already been given (for different other models) in [6]. Assume that $\{s_i\}_{i=1}^m$ is a partition for $s \in [s_{\min}, s_{\max}]$. Then, we define:

$$s_i = \Psi(\xi_i), \quad 1 \leq i \leq m, \quad (4)$$

wherein $m \gg 1$ and $\xi_{\min} = \xi_1 < \xi_2 < \dots < \xi_m = \xi_{\max}$ are m equally-spaced points with the following features: $\xi_{\min} = \sinh^{-1}\left(\frac{s_{\min} - s_{\text{left}}}{d_1}\right)$, $\xi_{\text{int}} = \frac{s_{\text{right}} - s_{\text{left}}}{d_1}$, $\xi_{\max} = \xi_{\text{int}} + \sinh^{-1}\left(\frac{s_{\max} - s_{\text{right}}}{d_1}\right)$. We also consider here that $s_{\min} = 0$. The parameter $d_1 > 0$ controls the density of the points $s = E$. Moreover, we have

$$\Psi(\xi) = \begin{cases} s_{\text{left}} + d_1 \sinh(\xi), & \xi_{\min} \leq \xi < 0, \\ s_{\text{left}} + d_1 \xi, & 0 \leq \xi \leq \xi_{\text{int}}, \\ s_{\text{right}} + d_1 \sinh(\xi - \xi_{\text{int}}), & \xi_{\text{int}} < \xi \leq \xi_{\max}. \end{cases} \quad (5)$$

Some common choices are $d_1 = \frac{E}{4}$, while $s_{\text{left}} = \max\{0.5, e^{-0.0025T}\} \times E$, $s_{\text{right}} = E$, $[s_{\text{left}}, s_{\text{right}}] \subset [0, s_{\max}]$, and $s_{\max} = 4E$. Similarly, if $\{v_j\}_{j=1}^n$ be a set of nodes, then we define:

$$v_j = d_2 \sinh(\varsigma_j), \quad 1 \leq j \leq n, \quad (6)$$

where $n \gg 1$ and $d_2 > 0$ controls the mesh concentration near $v = 0$. A common choice is [6]: $d_2 = \frac{v_{\max}}{200}$, where $v_{\max} = 3$. In addition, ς_j are equally-spaced nodes given by $\varsigma_j = (j - 1)\Delta\varsigma$, $\Delta\varsigma = \frac{1}{n-1} \sinh^{-1}\left(\frac{v_{\max}}{d_2}\right)$, for any $1 \leq j \leq n$.

4 RBF-FD estimates

For the case of the 1st derivative and without loss of generality, we consider the three unstructured points [8], $\{x_i - h, x_i, x_i + \omega h\}$, $\omega > 0, h > 0$, and write:

$$f'(x_i) \simeq \Xi_{i-1} f(x_{i-1}) + \Xi_i f(x_i) + \Xi_{i+1} f(x_{i+1}) = \hat{f}'(x_i), \quad (7)$$

wherein f and $\hat{f}$ stand for the exact and approximate values. As long as $c \gg h$, we attain

$$\begin{aligned}\Xi_{i-1} &= -\frac{\psi_1}{2c^2h(\omega+1)(-3(c^4-2)\theta^2+(c^6+3c^4-2)\theta^3-6\theta+2)}, \\ \Xi_i &= \frac{\psi_2}{2c^2h\omega(-3(c^4-2)\theta^2+(c^6+3c^4-2)\theta^3-6\theta+2)}, \\ \Xi_{i+1} &= -\frac{\psi_3}{2c^2h\omega(\omega+1)(-3(c^4-2)\theta^2+(c^6+3c^4-2)\theta^3-6\theta+2)},\end{aligned}\tag{8}$$

wherein $\psi_1 = \omega(2c^8\theta^3 + c^6\theta^2(\theta(h^2(4-3\omega)+6)-6) - 5c^4h^2(\theta-1)\theta^2(3\omega-4) + c^2(\theta-1)^2(\theta(h^2(9\omega-5)-4)+4) + h^2(\theta-1)^3(11\omega-17))$, $\psi_2 = (\omega-1)(2c^8\theta^3 - 3c^6\theta^2(\theta(h^2\omega-2)+2) - 15c^4h^2(\theta-1)\theta^2\omega + c^2(\theta-1)^2(\theta(9h^2\omega-4)+4) + 11h^2(\theta-1)^3\omega)$, and $\psi_3 = -2c^8\theta^3 + c^6\theta^2(\theta(h^2(3-4\omega)\omega-6)+6) - 5c^4h^2(\theta-1)\theta^2\omega(4\omega-3) + c^2(\theta-1)^2(\theta(h^2\omega(5\omega-9)+4)-4) + h^2(\theta-1)^3\omega(17\omega-11)$.

To calculate the weights for the function 2nd derivative, one obtains as long as $c \gg h$:

$$f''(x_i) \simeq \Theta_{i-1}f(x_{i-1}) + \Theta_i f(x_i) + \Theta_{i+1}f(x_{i+1}) = \hat{f}''(x_i),\tag{9}$$

where

$$\begin{aligned}\Theta_{i-1} &= \frac{\psi_4}{c^2h^2(\omega+1)(\theta(\theta(c^6\theta+3c^4(\theta-1)-2\theta+6)-6)+2)}, \\ \Theta_i &= -\frac{\psi_5}{c^2h^2\omega(\theta(\theta(c^6\theta+3c^4(\theta-1)-2\theta+6)-6)+2)}, \\ \Theta_{i+1} &= \frac{\psi_6}{c^2h^2\omega(\omega+1)(\theta(\theta(c^6\theta+3c^4(\theta-1)-2\theta+6)-6)+2)},\end{aligned}\tag{10}$$

where $\psi_4 = 2c^8\theta^3 + c^6\theta^2(\theta(h^2(\omega(3\omega-5)+4)+6)-6) + 5c^4h^2(\theta-1)\theta^2(\omega(3\omega-5)+4) - c^2(\theta-1)^2(\theta(h^2(\omega(9\omega-13)+5)+4)-4) - h^2(\theta-1)^3(\omega(11\omega-19)+17)$, $\psi_5 = 2c^8\theta^3 + c^6\theta^2(\theta(h^2(\omega(3\omega-4)+3)+6)-6) + 5c^4h^2(\theta-1)\theta^2(\omega(3\omega-4)+3) - c^2(\theta-1)^2(\theta(h^2(\omega(9\omega-17)+9)+4)-4) + h^2(\theta-1)^3((13-11\omega)\omega-11)$ and $\psi_6 = 2c^8\theta^3 + c^6\theta^2(\theta(h^2(\omega(4\omega-5)+3)+6)-6) + 5c^4h^2(\theta-1)\theta^2(\omega(4\omega-5)+3) - c^2(\theta-1)^2(\theta(h^2(\omega(5\omega-13)+9)+4)-4) - h^2(\theta-1)^3(\omega(17\omega-19)+11)$.

5 Double integral and final system

We discretize the multiple integral operator of (2) given by:

$$\mathcal{L}_I u(s, v, \tau) = \int_0^\infty \int_0^\infty u(sz^s, v + z^v, \tau) \mathbf{p}(z^s, z^v) dz^v dz^s.\tag{11}$$

It is necessary to first impose a transformation as follows:

$$z_1 = sz^s, \quad z_2 = v + z^v.\tag{12}$$

The transformation (12) acts on (11) and leads to:

$$\mathcal{L}_I u(s, v, \tau) = \int_0^\infty \int_v^\infty \frac{1}{s} u(z_1, z_2, \tau) \mathbf{p}(z_1/s, z_2 - v) dz_2 dz_1.\tag{13}$$

Computing (13) at the computational node (s_i, v_j) reads $\mathcal{L}_I u(s_i, v_j, \tau) = \int_0^\infty \int_{v_j}^\infty \frac{1}{s_i} u(z_1, z_2, \tau) \mathbf{p}(z_1/s_i, z_2 - v_j) dz_2 dz_1$. Here, we choose a rectangular $[0, s_{\max}] \times [v_j, v_{\max}]$, which is *large enough* so that integrating over it gives a sufficiently good approximation as in the work [7]. Accordingly, in the computational domain, we have:

$$\mathcal{L}_I u(s_i, v_j, \tau) = \int_0^{s_{\max}} \int_{v_j}^{v_{\max}} \frac{1}{s_i} u(z_1, z_2, \tau) \mathbf{p}(z_1/s_i, z_2 - v_j) dz_2 dz_1.\tag{14}$$

This yields $\mathcal{L}_I u(s_i, v_j, \tau) = \sum_{p=1}^{p=m-1} \sum_{q=j}^{q=n-1} M_{p,q}$, where

$$M_{p,q} = \int_{s_p}^{s_{p+1}} \int_{v_q}^{v_{q+1}} \frac{1}{s_i} u(z_1, z_2, \tau) \mathfrak{p} \left(\frac{z_1}{s_i}, z_2 - v_j \right) dz_2 dz_1. \quad (15)$$

We consider an estimate for u for every cell $[s_p, s_{p+1}] \times [v_q, v_{q+1}]$, and then integrate the density function $\mathfrak{p}$ theoretically. The simplest way to approximate u is to take into account u as a fixed function on each cell as comes next:

$$u(z_1, z_2, \tau) \simeq \frac{1}{4} (u(s_p, v_q, \tau) + u(s_p, v_{q+1}, \tau) + u(s_{p+1}, v_q, \tau) + u(s_{p+1}, v_{q+1}, \tau)), \quad (16)$$

where $(z_1, z_2) \in [s_p, s_{p+1}] \times [v_q, v_{q+1}]$. Using this approximation for u in (15) we have:

$$\begin{aligned} M_{p,q} = & \frac{1}{4} (u(s_p, v_q, \tau) + u(s_p, v_{q+1}, \tau) + u(s_{p+1}, v_q, \tau) + u(s_{p+1}, v_{q+1}, \tau)) \\ & \times \mathfrak{P} \left(\frac{s_p}{s_i}, \frac{s_{p+1}}{s_i}, v_q - v_j, v_{q+1} - v_j \right), \end{aligned} \quad (17)$$

where $\mathfrak{P}(A, B, C, D) = \int_A^B \int_C^D \mathfrak{p}(z_1, z_2) dz_2 dz_1$.

Ultimately after the incorporation of the boundaries directly into the system matrix, we contribute a set of semi-discretized (linear) ODEs with the system matrix $\tilde{\mathbf{Y}}$ which is real, and unsymmetric. The set of ODEs is locally well-posed, i.e., there exists a unique solution (depending on the initial condition), which satisfies the system with the Lipschitz constant $\|\tilde{\mathbf{Y}}\|$, for a spectral matrix norm.

This system is solved now by the Krylov subspace method which saves the computational time when the system matrix $\tilde{\mathbf{Y}}$ is of large size in contrast to the existing time-stepping solvers, at which the explicit ones relies heavily on the choice of a very refined step size and the implicit ones suffer from more computational burden imposed because of solving nonlinear (systems) of algebraic equations per temporal cycle.

References

- [1] L.V. Ballestra, C. Sgarra, *The evaluation of American options in a stochastic volatility model with jumps: an efficient finite element approach*, Comput. Math. Appl. 60 (2010), 1571–1590.
- [2] D. Bates, *Jumps and stochastic volatility: the exchange rate processes implicit in Deutsche mark options*, Rev. Fin. Studies, 9 (1996), 69–107.
- [3] D. Duffie, J. Pan, K. Singleton, *Transform analysis and asset pricing for affine jump diffusions*, Econometrica, 68 (2000), 1343–1376.
- [4] B. Eraker, M. Johannes, N. Polson, *The impact of jumps in volatility and returns*, J. Finance, 58 (2003), 1269–1300.
- [5] L. Feng, V. Linetsky, *Pricing options in jump-diffusion models: An extrapolation approach*, Oper. Res., 56 (2008), 304–325.
- [6] K.J. in 't Hout, S. Foulon, *ADI finite difference schemes for option pricing in the Heston model with correlation*, Int. J. Numer. Anal. Modeling, 7 (2010), 303–320.
- [7] S. Salmi, J. Toivanen, L. von Sydow, *An IMEX-scheme for pricing options under stochastic volatility models with jumps*, SIAM J. Sci. Comput., 36 (2014), 817–834.
- [8] I. Tolstykh, *On using RBF-based differencing formulas for unstructured and mixed structured-unstructured grid calculations*, Proc. 16th IMACS World Congress, 228 (2000), 4606–4624.

Email: fazlollah.soleymani@gmail.com

Pricing Arithmetic Asian option using the control variate technique

Davood Ahmadian¹

Department of Applied Mathematics, University of Tabriz, Tabriz, Iran

Abstract

Our goal is to develop new Monte Carlo techniques for pricing Asian options. In particular we consider discrete Asian arithmetic options driven by the mixed fractional Brownian motions and we price them using a control variate approach that exploits an explicit exact solution for discrete geometric Asian options.

Email: d.ahmadian1985@gmail.com

¹speaker

Valuing Equity-Linked Death Benefits under Phase-type models

Saghar Heidari²

Shahid Beheshti University, Tehran, Iran

Abstract

In this paper, we study the pricing problem of equity-linked life insurance products such as guaranteed minimum death benefit (GMDB) in the case that remaining life time of a policyholder, denoted τ , approximated by a Phase-type (PH) distribution and underlying asset dynamic, denoted by X_t , is described by a jump-diffusion regime-switching model. To find the fair value of the products, we use the discounted density approach. For this purpose, we modified the recent results for the Wiener-Hopf factorization under PH assumption in Markov-modulated economy to define the joint law of the overall supremum value of the process X_t at time τ .

Keywords: Life Insurance, Equity-Linked Death Benefits, Phase-type distribution, Regime Switching model

Mathematics Subject Classification [2018]: 13D45, 39B42

1 Introduction

In the past, life insurance companies used to offer life insurance contracts that provided a fixed capital at time of death of the policyholders embedded insurance protections. In the past decades, contracts often involve a return linked to the market. As an example of these contracts, equity-linked contracts are more popular products that offer minimum death benefits. The pricing problem of these life insurance products is a significant challenge for insurance companies for financial purposes.

It is known that the time-until-death random variable $\tau(x)$ of a person at age x , when signing the life insurance contract, is assumed as combination of exponential distributions [4]. We know that linear combination of exponential distributions are not necessary exponentially distributed and can be negative in the tail part, which plays an important role in valuing equity-linked insurance products. These difficulties leads to using alternative distributions such as PH distributions [5]. Approximation of the future lifetime distribution by PH instead of exponential distribution has two advantages. First is that many calculations that are explicit for exponential distributions are often computationally tractable with PH assumptions (Asmussen et al. 1996). Second reason is that, PH distributions are dense so that a given distribution F on $[0, \infty)$ can be approximated arbitrarily well by a PH distribution with large enough number of phases. There are many papers in ruin theory, credibility theory, risk theory and options pricing that considered PH distributions. But as far as we know, the only references of applications of PH distributions to life insurance are (Lin and Liu 2007) and (Zadeh et al. 2014).

Resend empirical studies show that the assuming Geometrical Brownian motion for dynamics of underlying assets in well-known Black-Scholes-Merton model (Black and Scholes 1972) has some deficiencies such as ignoring market jumps and considering market volatility as a constant factor which is not consistent with some financial features such as the volatility smile and the thin-tailed return distribution of asset

²speaker

dynamic. Consequently, alternative models have been introduced to capture those phenomena in financial markets. These models include stochastic volatility models, jump–diffusion models [3] and regime–switching models [2]. Applying these models provide better results on fitting market data because they can explain the jump patterns exhibited by some financial assets and have the potential to capture a wide variety of implied volatility skews in real world options data. In this framework, recent research in finance literature are allocated to creating more realistic models by combining different models like jump–diffusion plus regime–switching models. Under this model the parameters such as drift and volatility are allowed to take diverse values in finite number of regimes.

In this paper, we consider PH distributions for future life time of insured to find the fair value of equity-linked products in life insurance, an area where PH distributions have been far less employed than in the related areas of non-life insurance. We also assume that the underlying asset dynamics provided jump features of market price, considering jump diffusion regime switching models. To solve the pricing problem with discounted density approach, we apply a version of the Wiener-Hopf factorization. Recently the traditional version of Wiener-Hopf factorization has been developed [1] for the case of PH distribution of τ and underlying asset dynamic jump-diffusion with PH jumps. This motivated us to extend the results for the case of jump-diffusion regime-switching models.

The outline of the paper is as follows: The next section is allocated to introducing PH distributions. In section 3 we define the jump-diffusion regime-switching models for the dynamic of underlying assets. Then in section 4 we modified the Wiener-Hopf factorization for our model to find the joint law of the processes for pricing GMDB by discounted density approach. Finally the performance of the proposed model are illustrated through some numerical examples.

2 Phase-type Distributions

A random variable τ is called a phase-type variable with representation $(\alpha, \mathbf{T}, m)$, if $\tau = \inf\{t \geq 0, J_t = \dagger\}$ is the distribution of the life-time of a terminating time-homogeneous Markov process $\{J_t, t \geq 0\}$ on a finite state space with m states and additional absorbing state $\dagger$, where $\alpha = (\alpha_1, \dots, \alpha_m)$ is the vector of initial probability distribution and $\mathbf{T} = (t_{ij})_{m \times m}$ is generator matrix:

$$\begin{bmatrix} t_{11} & \cdots & t_{1m} \\ \vdots & \vdots & \vdots \\ t_{m1} & \cdots & t_{mm} \end{bmatrix}.$$

The associated vector of exit rates (killing rates) is given by $\mathbf{t} = -\mathbf{T}\mathbf{1}$, where $\mathbf{1}$ is a column vector with all entries equal to 1 and

$$\mathbf{t} = \begin{bmatrix} t_1 \\ \vdots \\ t_m \end{bmatrix},$$

with

$$t_i + \sum_{j=1}^m t_{ij} = 0, \quad i = 1, \dots, m.$$

The density function of the phase-type variable τ is given by

$$\mathbb{P}(\tau \leq z) = \alpha e^{\mathbf{T}z} \mathbf{t},$$

where the exponential of matrix $\mathbf{T}$ is defined as $\sum_{n=0}^{\infty} \frac{\mathbf{T}^n}{n!}$.

Assuming that $\mathbf{T} + \alpha\mathbf{t}$ is an irreducible matrix, because otherwise we can eliminate some states without affecting the distribution of τ , leads to this fact that all the non-terminating states are transient and so $\mathbf{T}$ is invertible.

We represent the time-reversed representation of τ by $(\alpha^*, \mathbf{T}^*, m)$ of terminating time-homogenous Markov process $\{J_t^*, t \geq 0\}$ with

$$J_t^* = \begin{cases} J_{(\tau-t)-}, & t < \tau, \\ \dagger, & \text{otherwise,} \end{cases}$$

where the initial distribution α^* , generator $\mathbf{T}^*$ and the exit rates vector of the process are as follows:

$$\alpha^* = \mathbf{t}^T \Delta_\nu, \quad \mathbf{T}^* = \Delta_\nu^{-1} \mathbf{T}^T \Delta_\nu, \quad \mathbf{t}^* = \Delta_\nu^{-1} \alpha^T,$$

where Δ_ν is the diagonal matrix with the positive vector $\nu = -\alpha \mathbf{T}^{-1}$ on the diagonal.

For more details of phase-type distributions, see (Asmussen 1996).

In this paper, we use phase-type distribution to approximate τ_x , the remaining life-time of an insured of age x , instead of exponential distribution.

3 Jump-Diffusion Regime-Switching model

We assume that the economy switches from one state to the other states governed by a finite state continuous time Markov chain $\{I(t)\}_{t \geq 0}$ with the state space $\{1, \dots, n\}$. Let the matrix $Q = (q_{i,j})_{n \times n}$ be the rate matrix of the chain. Recall that $q_{i,j} > 0$ for $i \neq j$ and that

$$\sum_{j=1}^n q_{i,j} = 0, \quad i = 1, \dots, n.$$

When the economy is in the i -th state (i.e., $I(t) = i$), the asset price process $S(t)$ is assumed to follow the jump-diffusion regime switching model:

$$\frac{dS(t)}{S(t-)} = \mu_i dt + \sigma_i dW(t) + dZ_i(t), \quad t > 0,$$

here, $\mu_i = r_i - d_i$, r_i is the interest rate, d_i is the dividend rate, σ_i is the volatility of the underlying asset, and $Z_i(t)$ is a compensated compound Poisson process, independent of the Wiener process $W(t)$ under the measure $\mathbb{Q}$. As usual, the process $Z_i(t)$ is specified as

$$Z_i(t) = \sum_{j=1}^{N_i(t)} (e^{J_{i,j}} - 1) - \lambda_i \epsilon_i t,$$

where ϵ_i is the expected jump percentage, $N_i(t)$ is a Poisson process with the risk-neutral intensity λ_i and $\{J_{i,j}\}_{j=1}^\infty$ is a sequence of independent and identically distributed random variables. For the two well-known jump models, Merton's model and Kou's model, the probability density functions of the random variable J_i are specified as follows:

$$g_i(y) = \frac{1}{\sqrt{2\pi}\sigma_{J_i}} e^{-\frac{(y-\mu_{J_i})^2}{2\sigma_{J_i}^2}}, \quad (1)$$

and

$$g_i(y) = p_i \eta_{i,1} e^{-\eta_{i,1} y} \mathbb{1}_{\{y \geq 0\}} + (1 - p_i) \eta_{i,2} e^{\eta_{i,2} y} \mathbb{1}_{\{y < 0\}}, \quad (2)$$

respectively, where $\mu_{J_i} \in \mathbb{R}$, $\sigma_{J_i} > 0$, $\eta_{i,1} > 1$, $\eta_{i,2} > 0$, $p_i \in [0, 1]$ and $\mathbb{1}$ is indicator function.

For each state define the Levy process $X_i(t)$ by asset price $S(t)$ such as $S(t) = S(0)e^{X_i(t)}$. From the Levy processes features (Cont 2004) the Levy-Khinchine formula $X_i(t)$ for complex values of w is given by

$$\begin{aligned} \Phi_{X_i(t)}(w) &= \mathbb{E}[\exp(wX_i(t))] \\ &= \exp t \left(\mu_i w + \frac{1}{2} \sigma_i^2 w^2 + \int_{\mathbb{R}} (e^{wx} - 1 - wx \mathbb{1}_{|x| \leq 1}) \lambda_i d\nu(x) \right) \\ &= \exp(t\kappa_i(w)), \end{aligned}$$

where

$$\kappa_i(w) = \mu_i w + \frac{1}{2} \sigma_i^2 w^2 + \lambda_i (\zeta_i - 1),$$

with

$$\begin{aligned} \zeta_i &= \frac{p_i \eta_{j1}}{\eta_{j1} - w} + \frac{(1 - p_i) \eta_{j2}}{\eta_{j2} + w}, \\ \zeta_i &= \exp(\mu_i w + \frac{1}{2} \sigma_i^2 w^2), \end{aligned}$$

for the Kou's model and Merton's model respectively.

Denote the elements $\mathbb{E}[\exp(wX_j(t))|I(0) = i]$ for $i, j = 1, \dots, n$, by matrix $[F_t(w)]_{n \times n}$ such that

$$F_t(w) = \exp(t\bar{\Phi}(w)),$$

with

$$\bar{\Phi}(w) = Q + A,$$

where A is diagonal matrix in which $A_{ii} = \kappa_i(w)$.

In this paper, we assume the dynamic of asset price as jump diffusion regime switching model and since our approach can be used by both jump models, we consider $g_i(y)$ as general case for jumps distribution.

4 Pricing Equity-Linked Life Insurance Products

In this section, we are interested to study some valuation problems of equity-linked products, in particular, guaranteed minimum death benefits (GMDBs) in life insurance.

We assume that the stock price process $S(t)$ follows the jump diffusion regime switching models as discussed in previous section. Let $\tau = \tau(x)$ be the time until death of an insured (policyholder) of age x at $t = 0$, when signing the contract. We propose that τ is phase-type distributed, i.e. the distribution in section 2, independent of $S(t)$.

In life insurance contracts, the payment of insurance benefits can be considered as general benefit function $\Phi(S(\tau), \bar{S}(\tau))$, where $\bar{S}(\tau)$ is the running maximum, $\bar{S}(\tau) = \max_{t \leq \tau} S(t)$. For example the benefit function of guaranteed minimum death benefits in life insurance contracts is

$$\Phi(S(\tau), \bar{S}(\tau)) = \max(S(\tau), K) = S(\tau) + (K - S(\tau))^+,$$

where K is the guaranteed amount and $a^+ = \max(a, 0)$.

So the guaranteed minimum death benefits can be considered as the payoff a Put vanilla option written on the stock.

To value the life insurance contracts, we need to calculate the expectation of discounted value of the benefits:

$$\mathbb{E}[e^{-r\tau} \Phi(S(\tau), \bar{S}(\tau))], \quad (3)$$

where r is a discounting factor.

Let the Levy process X_t as $S(t) = S(0)e^{X_t}$. Further, where J_t lives on phases $1, \dots, m$, X_t switched in regimes with $m(n+1)$ phases i, j for $i = 1, \dots, m$ and $j = 1, \dots, n$. The killing rate in state i is t_i and zero in all ij states.

To calculate the expectation of (3), the distribution of the maximum or minimum values of the X_τ is required, so we define

$$\begin{aligned} \bar{X}_t &= \sup_{s \leq t} X_s, \\ \underline{X}_t &= \inf_{s \leq t} X_s, \end{aligned}$$

and the times the maximum or minimum values are attained by

$$\begin{aligned}\bar{\eta}_t &= \inf\{s \leq t, X_s \vee X_{s-} = \bar{X}_t\}, \\ \underline{\eta}_t &= \sup\{s \leq t, X_s \wedge X_{s-} = \underline{X}_t\}.\end{aligned}$$

Now considering the reversal of PH, we have the following result.

Theorem 4.1. *Assumption the standard reversal of PH with representation $(\alpha^*, \mathbf{T}^*, m)$, the modified factorization is as follow*

$$\begin{aligned}\mathbb{P}_i(\bar{X}_\tau \in dx, X_\tau - \bar{X}_\tau \in dy, I_{\bar{\eta}_\tau} = k, J_{\tau-} = j) \\ = \mathbb{P}_i(\bar{X}_\tau \in dx, I_{\bar{\eta}_\tau} = k) \mathbb{P}_j^*(\underline{X}_\tau \in dy, I_{\underline{\eta}_\tau} = k) u_k,\end{aligned}$$

where,

$$u_k = \frac{\alpha_j^*}{\mathbb{P}(I_{\bar{\eta}_\tau} = k)} = \frac{\alpha_j^*}{\mathbb{P}^*(I_{\underline{\eta}_\tau} = k)}.$$

References

- [1] P.J.; Yang H. Asmussen, S.; Laub. Phase-type models in life insurance: Fitting and valuation of equity-linked benefits. *Risks*, 7, 2019.
- [2] J. Buffington and R. J. Elliott. American options with regime switching. *Int. J. Theor. Appl. Finance*, 5:497–514, 2002.
- [3] R. Cont and P. Tankov. *Financial Modeling With Jump Processes*. Chapman and Hall CRC Press, 2004.
- [4] Elias S.W. Shiu Gerber, Hans U. and Hailiang Yang. Valuing equity-linked death benefits in jump diffusion models. *Insurance: Mathematics and Economics*, 53:615–623, 2013.
- [5] O. Nerman S. Asmussen and M. Olsson. Fitting phase-type distributions via the em algorithm. *Scand. J. Statist.*, 23:419–441, 1995.

Email: `s_heidari@sbu.ac.ir`

A machine learning approach to parametrize implied volatility in a risk management problem

Hirbod Assa

University of Liverpool, Liverpool, United Kingdom

Mostafa Pouralizadeh¹

Allameh Tabataba'i University, Tehran, Iran

Abstract

The main goal of this paper is to price a risk management problem based on the price of European call options. First of all, we provide a risk management problem which plays the role of an insurance contract to manage the risk of losses caused by market price fluctuation. Then, working towards controlling price movements, we introduce a machine learning algorithm with a quadratic hypothesis to model implied volatility and rule out static arbitrage on call option surface. We address how to preclude over-fitting (high variance) and under-fitting (high bias) by a regularized cost function including a penalty term which controls the trade-off between over-fitting and under-fitting. Eventually, the results of a numerical implementation show that the proposed modeling of implied volatility yield a volatility surface free of static arbitrage, therefore it can be used to monitor price variability and also to improve the precision of contract pricing.

Keywords: Risk Management; Implied Volatility; Static Arbitrage; Machine Learning.

AMS Mathematical Subject Classification [2018]: 97M30

1 Main results

1.1 Risk management framework

let $(S_t)_{0 \leq t \leq T}$ be a random process representing the value of a risky asset over the interval $[0, T]$, and let $r > 0$ denote the risk-free interest rate. If an agent invests in the asset at time 0, we consider a problem that the agent wants to manage the risk of losses $L = (S_0 - e^{-rT}S_T)^+$. To manage the risk of losses the agent buys a contract X at the price $\pi(X)$, that also satisfies $0 \leq X \leq L$. Therefore, the agent's global position of risk then is $P = L - X + \pi(X)$. Also, to avoid the risk of moral hazard we assume X and $L - X$ are comonotonic; in other words, both parties in a contract feel the risk of losses. In a complete market we can assume the valuation π of the contract X is given by $\pi(X) = E^Q(X)$. To price a risk management contract, the most common approach is to minimize the agent's global position of risk. On the other hand, we assume that the agent measures the risk of losses with value-at-risk (VaR) introduced

¹speaker

as $\text{VaR}_\alpha(L) = \inf \{x \in R \mid F_X(x) \geq \alpha\}$, where $\alpha \in (0, 1)$ represents the risk aversion of the agent. Then, the risk management problem is given by:

$$\min_{\substack{0 \leq X \leq L, \\ X \text{ and } L-X \\ \text{are comonoton}}} \{ \text{VaR}_\alpha(L - X) + E^Q(X) \} \quad (1)$$

Now, based on [1], we provide the following theorem that links the price of risk management contracts to the price of European call options.

Theorem 1.1. *An optimal solution X for the problem 1 is given by*

$$X = \min \left\{ (S_0 - e^{-rT} S_T)^+, (S_0 - e^{-rT} \text{VaR}_{(1-\alpha)}(S_T))^+ \right\}$$

The form of the solution is given as $X = f(S_T)$, where

$$\begin{aligned} f(x) = & (S_0 - e^{-rT} \text{VaR}_{1-\alpha}(S_T))^+ + e^{-rT} (x - S_0 e^{-rT})^+ \\ & - e^{-rT} (x - \text{VaR}_{1-\alpha}(S_T))^+ \end{aligned} \quad (2)$$

In fact, an expectation over the solution X comes up with this idea that the first term of the solution 2 is a constant which can easily be computed by market information, but the second and third terms are somehow similar to the pricing formula for an European call option based on the Black-Scholes framework, so the problem of contract pricing is changed into the problem of option pricing, and we only need to make sure that the call options are correctly priced. This will lead us to restrict our analysis only to static arbitrage since it is a kind of arbitrage defined on the call surface and without loss of generality we consider that there is no martingale measure in the market. In the next step, we propose a quadratic machine learning approach to model the Black-Scholes implied volatility of European call option to take care of the contract pricing.

1.2 The quadratic parameterization

For a parameter set $\eta = \{\theta_0, \theta_1, \theta_2\}$, the quadratic parameterization of total implied variance with respect to moneyness x is given by $w_{imp}^{Q^2}(x, \eta) = \theta_0 + \theta_1 x + \theta_2 x^2$, where $\theta_0 > 0$, $\theta_1 \in \mathbb{R}$. The condition of $\theta_2 > 0$ along with the condition of $\theta_1^2 - 4\theta_0\theta_2 < 0$ make the function $x \rightarrow w_{imp}^{Q^2}(x, \eta)$ positive and strictly convex for all $x \in \mathbb{R}$.

Proposition 1.2. *The quadratic surface $w_{imp}^{Q^2}(x, \eta)$ is free of calendar spread arbitrage if for any two times to maturity $\tau_1 < \tau_2$ corresponding to $w(\cdot, \tau_1)$ and $w(\cdot, \tau_2)$ by the parameters sets $\eta_1 = \{\theta_{01}, \theta_{11}, \theta_{21}\}$ and $\eta_2 = \{\theta_{02}, \theta_{12}, \theta_{22}\}$ the following conditions satisfy*

- (1) $\theta_{22} - \theta_{21} > 0$;
- (2) $\theta_{22}\theta_{01} + \theta_{21}\theta_{02} < \frac{\theta_{12}\theta_{11}}{2}$.

Proposition 1.3. *The quadratic volatility model $w_{imp}^{Q^2}(x, \eta)$ for options with less than one year expiration time, is free of butterfly arbitrage if*

- (1) $\theta_1^2 - 4\theta_0\theta_2 < 0$;
- (2) $\frac{1}{4} < \theta_0 < 1$.

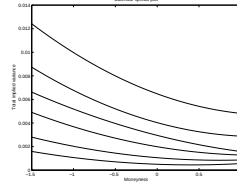


Figure 1: Calendar spread plot

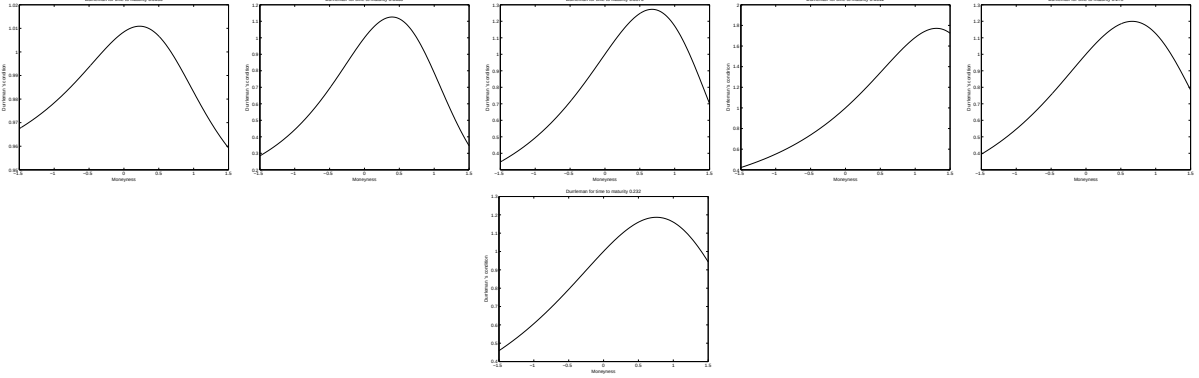


Figure 2: Durrleman Plots for six different expiration time.

1.3 The cost function

To avoid high bias (under-fitting) and high variance (over-fitting), we use a penalty term [4] for each volatility slice as follows

$$\hat{\eta} = \arg \min_{\theta} \frac{1}{2m} \left(\sum_{i=1}^m \left(w_{\theta}^{Q^2}(x^{(i)}) - w^{(i)} \right)^2 + \lambda \sum_{j=1}^2 \theta_j^2 \right) \quad (3)$$

and Also a forward slice by slice strategy is used in this paper to rule out calendar spread arbitrage. Figure 1 shows that the parametrization is free of calendar spread arbitrage since total implied variance is an increasing function of expiry time [3]. Figure 2 shows that for all different times to expiration, the volatility slice is free of butterfly arbitrage because the Durrleman's function for all volatility slices is positive around at the money [2]. Table 1 shows the information about time to maturity and the optimum value of λ for each volatility slice.

 Table 1: Expiration times and the optimum values of λ

Expiry date	Time to maturity	λ
20/12/2014	0.0136	0.3
02/01/2015	0.0465	3
17/01/2015	0.0876	1.2
23/01/2015	0.1041	2.8
20/02/2015	0.178	0.9
20/03/2015	0.232	1.3

Acknowledgment

The acknowledgements should be in a separate section at the end of the article before the references.

References

- [1] ASSA, H. *On Optimal Reinsurance Policy with Distortion Risk Measures and Premiums*, *Insurance. Mathematics and Economics* (2015), Volume 61, pp 70–75.
- [2] DURRLEMAN, V. *A note on initial volatility surface* (2003), Unpublished manuscript.
- [3] J. GATHERAL, *The volatility surface: a practitioner's guide*, Vol. 357, John Wiley & Sons (2011).
- [4] HASTIE, T., R. TIBSHIRANI, AND J. FRIEDMAN. (2002). *The Elements of Statistical Learning, Data Mining, Inference, and Prediction*. Vol. 1. New York: Springer series in statistics.
- [5] M. ROPER, *Implied volatility: General properties and asymptotics*, The University of New South Wales (2009), 2-3.

Lapse risk modeling with machine learning techniques: a case study on Mellat insurance policyholders

Khaled Masoumifard

Mellat Insurance, Tehran, Iran

Soroush Amirhashchi¹

Mellat Insurance, Tehran, Iran

Reyhaneh Jannati

Mellat Insurance, Tehran, Iran

Abstract

In this paper, Our objective is to predict the lapse risk by studying savings contracts at an individual level. In fact, by considering a sequence of a binary process for a policyholder status, we present an algorithm based on dynamic machine learning to predict its path. Using this method, we provide a dynamic model to predict the lapse risk of each policyholder based on static and dynamic factors.

Keywords: Lapse, Cash-Flow-Projection, Solvency Capital Requirement

AMS Mathematical Subject Classification [2018]: C52, C53, G22

1 Introduction

In this paper, we focus on structural lapse risk in the case of Individual savings contracts and apply machine learning methods to identify homogeneous risk groups. One of the interesting features of the methods used here is to be able to model the behavior of policyholders based on static and dynamic factors (such as economic factors). To study the impact of dynamic and static factors on lapse risk, see [1].

2 Model formulation

The purpose of this paper is to describe a sequence of machine learning procedures that can be characterize the process of lapse risk. In this approach, we first introduce a machine learning algorithm to train the parameters of each customer's behavior prediction model in the first year of its contract, and then deal with model training with another machine learning algorithm for each customer's behavior in the second

¹speaker

year of its contract. We continue this approach until that we can extract a measure of the behavior of each customer with respect to the lapse process. Consider the following data structure for further explanation

$$\begin{aligned}
&\text{Policyholder}_1 \longrightarrow (1, \mathbf{X}_{y_1}^{(1)}) \longrightarrow (1, \mathbf{X}_{y_1+1}^{(1)}) \longrightarrow (0, \mathbf{X}_{y_1+2}^{(1)}) \\
&\text{Policyholder}_2 \longrightarrow (1, \mathbf{X}_{y_2}^{(2)}) \longrightarrow (1, \mathbf{X}_{y_2+1}^{(2)}) \longrightarrow (1, \mathbf{X}_{y_2+2}^{(2)}) \longrightarrow \dots \\
&\text{Policyholder}_3 \longrightarrow (1, \mathbf{X}_{y_3}^{(3)}) \longrightarrow (1, \mathbf{X}_{y_3+1}^{(3)}) \longrightarrow (1, \mathbf{X}_{y_3+2}^{(3)}) \longrightarrow (1, \mathbf{X}_{y_3+3}^{(3)}) \longrightarrow (0, \mathbf{X}_{y_3+4}^{(3)}) \quad . \\
&\text{Policyholder}_4 \longrightarrow (0, \mathbf{X}_{y_4}^{(4)}) \\
&\text{Policyholder}_5 \longrightarrow (1, \mathbf{X}_{y_5}^{(5)}) \longrightarrow \dots
\end{aligned}$$

In general, investigate data with this structure is complicated: for example consider policyholder 2, which has only been observed for two years of its process, and we do not know what will happen to it in the future, but all situations of the policyholder 3 have been observed for 10 years. Therefore we have a lot of censored observations. To extract information about risk lapse from this data in an appropriate way, at first we remove all policyholders that dead and then classify these data into k sets. The structure of data in class i , $i \leq k$, as follows,

$$\begin{aligned}
&\text{Policyholder 1} \longrightarrow (1, \mathbf{X}_{y_1}^{(1)}) \longrightarrow (1, \mathbf{X}_{y_1+1}^{(1)}) \longrightarrow (1, \mathbf{X}_{y_1+2}^{(1)}) \longrightarrow \dots \longrightarrow (1, \mathbf{X}_{y_1+i}^{(1)}) \\
&\text{Policyholder 2} \longrightarrow (1, \mathbf{X}_{y_2}^{(2)}) \longrightarrow (1, \mathbf{X}_{y_2+1}^{(2)}) \longrightarrow (1, \mathbf{X}_{y_2+2}^{(2)}) \longrightarrow \dots \longrightarrow (1, \mathbf{X}_{y_2+i}^{(2)}) \\
&\quad \quad \quad \vdots \quad \quad \quad \vdots \quad \quad \quad \vdots \quad \quad \quad \vdots \quad \quad \quad \dots \quad \quad \quad \vdots \quad \quad \quad . \\
&\text{Policyholder } j \longrightarrow (1, \mathbf{X}_{y_j}^{(j)}) \longrightarrow (1, \mathbf{X}_{y_j+1}^{(j)}) \longrightarrow (1, \mathbf{X}_{y_j+2}^{(j)}) \longrightarrow \dots \longrightarrow (0, \mathbf{X}_{y_j+i}^{(j)}) \\
&\quad \quad \quad \vdots \quad \quad \quad \vdots \quad \quad \quad \vdots \quad \quad \quad \vdots \quad \quad \quad \dots \quad \quad \quad \vdots
\end{aligned}$$

This means that in class i , the probability of Lapse is investigated for a person who at least four years of their contract elapsed. Therefore, all people who have been under contract for less than i years or have been lapsed before i years of their contract elapsed are removed from class i and then we classify the remaining data using a supervised machine learning algorithm.

3 Main result

In this paper, we investigate the lapse risk in life insurance of Mellat Insurance Company by applying the model described in the previous section in life insurance data from 2009 to 2019. We have applied the 5-fold cross validation method to classification this data in all ten steps of the machine learning algorithm. At each step, the response variable is a binary variable such that one represents lapse at that step and zero indicates the contract is active or ended. Also, the age of policyholder and inflation rate were used as independent variables in this model. Implementation results show that the lapse risk can be predicted with excellent accuracy in the future years. The results for model accuracy are presented in Figure 1 and Table 1. Thus, with these results, it can be claimed that the model described in the previous section can predict the behavior of each policyholder with respect to attributes of each policyholder, economic conditions and etc.

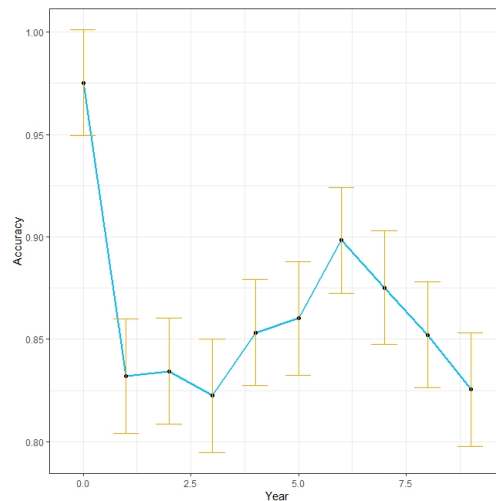


Figure 1: Accuracy for lapse projection

Iteration	Accuracy	Kappa	AccuracySD	KappaSD
Lapse in arrival year	0.9753	0.3320	0.0003	0.0089
Lapse after one year	0.8284	0.3042	0.0029	0.0126
Lapse after two year	0.8316	0.5686	0.0060	0.0171
Lapse after three year	0.8180	0.5184	0.0022	0.0058
Lapse after four year	0.8494	0.6538	0.0031	0.0072
Lapse after five year	0.8521	0.7062	0.0094	0.0184
Lapse after six year	0.8932	0.7591	0.0132	0.0315
Lapse after seven year	0.8588	0.6205	0.0257	0.0621
Lapse after eight year	0.8327	0.5244	0.0090	0.0359
Lapse after nine year	0.8229	0.3859	0.0188	0.0793

Table 1: projection accuracy of the lapse risk model for 10 years

Acknowledgment

The authors gratefully acknowledge support from the Mellat insurance company.

References

- [1] F. Barsotti and X. Milhaud, and Y. Salhi, *Lapse risk in life insurance: Correlation and contagion effects among policyholders' behaviors*, Insurance: Mathematics and Economics, no. 71 (2016), pp. 317-331.

e-mail: k_masoumifard@sbu.ac.ir
e-mail: soroush.amirhashchi@gmail.com
e-mail: R.jannati.k@gmail.com

Option pricing under a class of time-changed Levy process

Navideh Modarresi
Allameh Tabataba'i University, Tehran, Iran
Mahdiye Alijani¹
Allameh Tabataba'i University, Tehran, Iran

Abstract

Some empirical evidence shows that Jumps effect on financial asset prices and volatilities of returns change randomly over the time. In this paper, we present a class of time-change Levy process that can capture the volatility clustering and heteroskedasticity. The random clock process used to induce stochastic volatility is the integral of a positive increasingly Levy process with mean reverting property. By using the characteristic function and fast Fourier transform technique, the closed form formula of option pricing can be derived. In addition, taking into account this structure, we show that random clocks driving are highly correlated to trading volumes. Finally, after calibrating the parameters of the process we perform empirical researches on two Iran stock indices.

Keywords: European option, Time-changed processes, Subordinator, Volatility clustering.

Mathematics Subject Classification [2010]: 60G51; 60H10; 60E10.

1. Introduction

One of the most important topics in financial markets is the pricing of securities that are used as a hedging instrument and the value of the securities is derived from value of their underlying asset. Option pricing is the most common field in study of derivatives market. Geometric Brownian motion is used to model option prices in the Black-Scholes model. Financial market data shows jumps in prices, the skewness in distribution of returns compared to the normal distribution and as well as stochastic volatility over time. Therefore, the prices of the option resulting from the Black-Scholes model are not consistent with market data. In order to enhancing the performance of the Black-Scholes model for option pricing, a model that covers price jumps in asset pricing models is used [4]. The value of an option has dependence to the current stock price, the time value, volatility and interest rates. So the underlying asset must represent characteristics of financial markets such as skewness, heavy tailedness, volatility clustering and jump in asset price. The volatility clustering appears when the volatility of their returns are correlated [1]. These features can be found in a family of Levy's processes that are more in line with empirical realities [2]. For this purpose, we can randomize the volatility parameter of the Black-Scholes model or to time-change Levy process with a positive increasing process with dependent increments to cover the stochastic volatility [3]. The volatility clustering feature indicates that asset returns are not independent across time while Levy processes have independent increments, seems to be weak in representing the volatility clustering. For this reason, the stochastic volatility models are applied to represent volatility clustering. One approach into discrete-time models is dealing with GARCH models. In continuous-time processes by randomizing the volatility parameter such as the Heston model whose volatility is the Cox-Ingersoll-Ross (CIR) process and also the Bates model we can introduce some kind of stochastic volatility processes [5]. Time-changed Levy processes offer an alternative

¹speaker

method to generate stochastic volatility with non-stationary increments consist to use an integrated mean reverting process like CIR [3]. In this article, we apply a class of time-changed Levy process subordinated by a positive increasingly Levy process for option pricing. This process has Markov property but is not time-homogeneous. One of the advantages of these processes is that there exist analytical expressions of the mean, variance and moment generating function (mgf) of them. The other one is that one can consider the corresponding multi-variable process as a linear combination of different processes, with the same random times, and obtain their stochastic correlation function. In this framework first we calibrate the parameters of the model and fit the process to daily log returns of two indices of Iran stock market. The numerical analysis shows that the introduced stochastic times are highly correlated to trading volumes.

2. Theoretical framework of the model

In this section, first by the famous Levy-khintchine formula we present the moment generating function (mgf) of the process and present the dynamic of the jump diffusion process. Then introduce it's subordinator. In the univariate case, let $(X_t)_{t \geq 0}$ be a Levy process defined on a filtered probability space denoted by $(\Omega, (\mathcal{G}_t)_{t \geq 0}, \mathbf{P})$. This process has independent and stationary increments described by a triplet (μ, σ^2, ν) , The mgf of the process is denoted by

$$\phi_t^X(\omega) := \mathbb{E}(\exp(\omega(X_t - X_0)) | \mathcal{G}_0) = \exp(t\psi_X(\omega)), \quad (1)$$

Where $\psi_X(\omega) = (\mu\omega + \frac{1}{2}\omega^2\sigma^2 + \int_{\mathbb{R}} (e^{\omega z} - 1 - \omega z 1_{|z| < 1})\nu(dz))$ is the characteristic exponent. Without loss of generality, we assume that $X_0 = 0$. According to the Levy-Ito decomposition, X_t is the sum of three components: a deterministic drift μt , a diffusion with variance σ^2 and a jump process, $J_X(t, z)$, of intensity $\nu(\cdot)$, also called the Levy measure of X_t . So the jump diffusion process can be represented as $dX_t = \mu dt + \sigma dW_t + \int_{|z| > 1} z J_X(dt, dz) + \int_{|z| \leq 1} z (J_X(dt, dz) - \nu(dz) dt)$ where W_t is a Brownian motion. The stochastic clock of the time-changed Levy process is the integral of a positive increasing process λ_s as $\tau_t := \int_0^t \lambda_s ds$. λ_t represents the arrival rate of information. This process is also the intensity of a jump process $L_t = \sum_{k=1}^{N_t} J_k$, where N_t is a Poisson process and J_k is a sequence of independent random jumps with an exponential density, $\delta(z) = \rho e^{-\rho z} 1_{(z \geq 0)}$. The stochastic differential equation of the λ_t is $d\lambda_t = \alpha(\theta - \lambda_t)dt + \eta dL_t$ where $\alpha \in \mathbb{R}^+$ is a speed of reversion toward $\theta \in \mathbb{R}^+$, a mean level which has the solution as

$$\lambda_t = \theta + e^{-\alpha t}(\lambda_0 - \theta) + \eta \int_0^t e^{-\alpha(t-s)} dL_s.$$

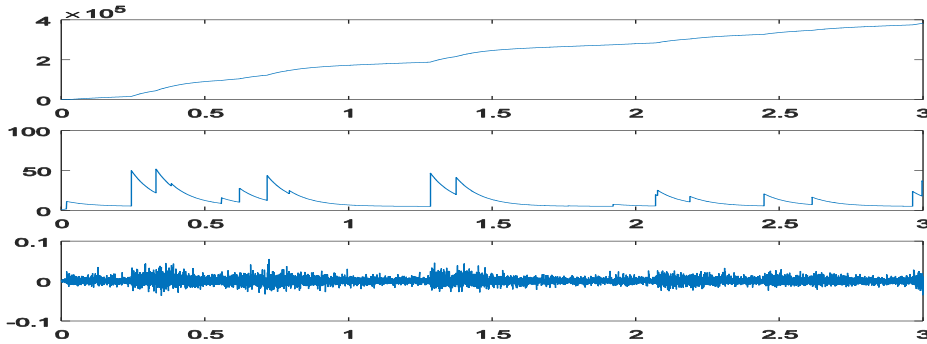


Fig.1. The first and second subplots show the sample paths of τ_t and λ_t , daily simulated over a period of 3 years and the third of is the daily variations of X_{τ_t} .

In Figure 1, we show the sample paths of τ_t , λ_t

and X_{τ_t} respectively. Parameters for this simulation are $\alpha = 11.5$, $\theta = 3$, $\eta = 1$ and $\rho = 0.1159$. X_t is a normal inverse Gaussian (NIG) process with a mean of 5% and a deviation of 15% .

To investigate the properties of the time-changed Levy process, the conditional mean, conditional variance and covariance function of the introduced process by using these moments are given. For more details See [6].

The expected value of X_{τ_t} , conditionally on $\mathcal{F}_0$, is given by

$$\mathbb{E}(X_{\tau_t} | \mathcal{F}_0) = \left(\mu + \int_{\mathbb{R}} z(1 - I_{|z| < 1})\nu(dz) \right) \mathbb{E}(\tau_t | \mathcal{F}_0) \quad (2)$$

where

$$\mathbb{E}(\tau_t | \mathcal{F}_0) = \frac{1}{\left(\frac{\eta}{\rho} - \alpha\right)} \left(\frac{\alpha\theta}{\frac{\eta}{\rho} - \alpha} + \lambda_0 \right) \left(e^{\left(\frac{\eta}{\rho} - \alpha\right)t} - 1 \right) - \frac{\alpha\theta}{\frac{\eta}{\rho} - \alpha} t. \quad (3)$$

Also, The variance of X_{τ_t} , conditionally to $\mathcal{F}_0$, is given by

$$\begin{aligned} \mathbb{V}(X_{\tau_t} | \mathcal{F}_0) &= \left(\sigma^2 + \int_{\mathbb{R}} z^2 v(dz) \right) \mathbb{E}(\tau_t | \mathcal{F}_0) + \mathbb{E}(X_1 | \mathcal{F}_0)^2 \mathbb{V}(\tau_t | \mathcal{F}_0) \\ &= \mathbb{V}(X_1 | \mathcal{F}_0) \mathbb{E}(\tau_t | \mathcal{F}_0) + \mathbb{E}(X_1 | \mathcal{F}_0)^2 \mathbb{V}(\tau_t | \mathcal{F}_0) \end{aligned} \quad (4)$$

If $X_{1,\tau_t}, X_{2,\tau_t}, \dots, X_{n,\tau_t}$ are subordinated by the same stochastic clock τ_t , the covariance between X_{i,τ_t} and X_{j,τ_t} for $i \neq j$, conditionally to $\mathcal{F}_0$, is time varying as

$$\mathbb{C}(X_{i,\tau_t}, X_{j,\tau_t} | \mathcal{F}_0) = \mathbb{E}(X_{i,1} | \mathcal{F}_0) \mathbb{E}(X_{j,1} | \mathcal{F}_0) \mathbb{V}(\tau_t | \mathcal{F}_0) \quad (5)$$

where $\mathbb{V}(\tau_t | \mathcal{F}_0)$ is defined by Eq. (4).

Furthermore, we can consider n time-changed Levy processes may be linearly combined to define a new multivariate process with a richer dynamics. For example, for the bivariate model ($n = 2$) is defined by linear combinations of X_{1,τ_t} and X_{2,τ_t} as $Y_{1,t} = X_{1,\tau_t}$ and $Y_{2,t} = mX_{1,\tau_t} + X_{2,\tau_t}$, where $m \in \mathbb{R}$ is a constant. In this framework, the time varying covariance between $Y_{1,t}$ and $Y_{2,t}$ is given

$$\mathbb{C}(Y_{1,t}, Y_{2,t} | \mathcal{F}_0) = m \mathbb{V}(X_{1,\tau_t} | \mathcal{F}_0) + \mathbb{E}(X_{1,1} | \mathcal{F}_0) \mathbb{E}(X_{2,1} | \mathcal{F}_0) \times \mathbb{V}(\tau_t | \mathcal{F}_0) \quad (6)$$

where $\mathbb{E}(\tau_t | \mathcal{F}_0)$ and $\mathbb{V}(\tau_t | \mathcal{F}_0)$ are defined by expressions (3) and (4). The mgf of X_{τ_t} is

$$\begin{aligned} \phi_t^{X_{\tau_t}}(\omega) &:= \mathbb{E}(\exp(\omega X_{\tau_t}) | \mathcal{F}_0) = \mathbb{E} \left(\mathbb{E} \left(\exp(\omega(X_{\tau_t})) \middle| \mathcal{F}_0 \vee \mathcal{H}_t \right) \middle| \mathcal{F}_0 \right) \\ &= \mathbb{E}(\exp(\bar{\omega} \tau_t) | \mathcal{F}_0) \end{aligned} \quad (7)$$

Where $\bar{\omega} := \left(\mu\omega + \frac{1}{2} \omega^2 \sigma^2 + \int_{\mathbb{R}} (e^{\omega z} - 1 - \omega z 1_{|z| < 1}) v(dz) \right)$ and $\mathcal{H}_t$ is the filtration of λ_{τ_t} .

3. Pricing of Options

The pricing of derivatives is done in a risk neutral world. And under the considered measure, discounted prices of assets are martingales to ensure the absence of arbitrage. It is then crucial to determine a family of equivalent measures, eligible to be risk neutral. Let us consider the family of equivalent measures that is induced by exponential martingales of the form:

$$M_t(Y, \xi) := \exp \left(g\lambda_{\tau_t} + YL_{\tau_t} - \varphi t + \xi X_{\tau_t} - \tau_t \psi_X(\xi) \right) \quad (8)$$

where Y, ξ, g and φ are constant. Is shown that M_t is a martingale, which is an essential feature to use M_t as a change of measure. Under the equivalent measure $\mathbb{Q}^{Y, \xi}$, X_t is a process with mgf $\phi_t^{X, \mathbb{Q}}(\omega) := \exp(t\psi_X^{\mathbb{Q}}(\omega) | \mathcal{F}_0)$

where $\psi_X^{\mathbb{Q}}(\omega) = \frac{1}{\phi^J(g\eta + Y)} (\psi_X(\xi + \omega) - \psi_X(\xi))$ and $N_t^{\mathbb{Q}}$ is a counting process with an intensity $\lambda_t^{\mathbb{Q}} = \phi^J(g\eta + Y) \lambda_t$. In order to invert numerically the probability density function $f_{X_{\tau_t}}(\cdot)$ to estimate the

parameters by log-likelihood maximization, we discretize the function by Discrete Fourier Transform (DFT). Let M be the number of steps used in the DFT method and $\Delta_X = \frac{2X_{\max}}{M-1}$, be the step of discretization. Let us denote

$$\Delta_Z = \frac{2\pi}{M\Delta_X} \text{ and}$$

$Z_j = (j-1)\Delta_Z$, for $j = 1, \dots, M$ and $\phi_T^{X_{\tau_t}}$ be mgf of X_{τ_t} , under the risk neutral measure $\mathbb{Q}$.

The values of $f_{X_{\tau_t}}(\cdot)$ at points $X_K = -\frac{M}{2}\Delta_X + (K-1)\Delta_X$ are computed approximately as

$$f_{X_{\tau_t}}(X_K) \approx \frac{2}{M\Delta_X} \operatorname{Re} \left(\sum_{j=1}^M \gamma_j \phi_T^{X_{\tau_t}}(iZ_j) (-1)^{j-1} e^{-i\frac{2\pi}{M}(j-1)(K-1)} \right) \quad (9)$$

where $\gamma_j = \frac{1}{2} I_{[j=1]} + I_{[j \neq 1]}$ and $\phi_T^{X_{\tau_t}}(iZ_j) = \mathbb{E}(e^{iZ_j X_{\tau_t}} | \mathcal{F}_0)$. Furthermore, we consider a sample of discrete observations of X_t as X_j for $j = 1, \dots, n$ where the interval of time between two observations is Δ . For filtering the market time a particle filter is proposed and in the structure of the algorithm, the log likelihood is approached. Next, by maximizing the log likelihood the parameters are estimated.

Let us consider an European option of maturity T , written on a single stock, with a price denoted $S_t = S_0 e^{X_{\tau_t}}$. The payoff is expressed as a function of its log-return and is noted $V(S_t)$. The option price under the risk neutral measure is

$$O(t, X_t, \lambda_t, \tau_t) = E^Q \left(e^{-\int_t^T r(s) ds} V(e^{X_{\tau_t}}) | \mathcal{F}_t \right) = E^Q \left(e^{-\int_t^T r(s) ds} V(X_{\tau_t}) | \mathcal{F}_t \right). \quad (10)$$

This numerical application aims to test the ability of time-changed Levy processes to explain the behavior of Basic metals sector indices and the Metal ore sector indices. Three time-changed Levy processes are considered in this exercise: a Brownian motion (BM), a normal inverse Gaussian (NIG) and a Variance Gamma (VG) process. For fitting the process, one simple approach is to consider the arrival rate of trades constant and equal to its asymptotic level $\lambda_\infty := \lim_{t \rightarrow \infty} E[\lambda_t] = -\frac{\alpha\theta}{\rho}$.

Under this assumption, the proposed model becomes stationary and, for a given set of parameters θ , daily returns have the same statistical distribution $f_{X_\tau}(\cdot | \theta, \lambda_\infty)$. Parameters are finally estimated by the following log-likelihood maximization:

$$\theta^{opt} = \arg \max_{\theta} \sum_{j=1}^T \log \left(f_{X_\tau}(y_j | \theta, \lambda_\infty) \right)$$

By the covariance function in Eq. (6), the correlation between the returns in the multivariate model is given by:

$$\text{cor}(Y_{1,t}, Y_{2,t} | \mathcal{F}_0) \approx \frac{mV(X_{1,\tau_t} | \mathcal{F}_0) + E(X_{1,1} | \mathcal{F}_0)E(X_{2,1} | \mathcal{F}_0) \times V(\tau_t | \mathcal{F}_0)}{\sqrt{V(X_{1,\tau_t} | \mathcal{F}_0)} \sqrt{m^2 V(X_{1,\tau_t} | \mathcal{F}_0) + V(X_{2,\tau_t} | \mathcal{F}_0)}}. \quad (11)$$

The comparison of standard deviations of X_{1,τ_t} and X_{2,τ_t} emphasizes that the Basic metals and the Metal ore are two stocks indices with very similar variances, then we infer that the correlation between daily returns is nearly constant and close to $\frac{m}{\sqrt{1+m^2}}$.

References

- [1] P. Carr, L. Wu, Time-changed Lévy processes and option pricing, Journal of Financial Economics. 71 (2004) 113–141.
- [2] R. Cont, P. Tankov, Financial modeling with jump process, CRC Press LIC 2004.
- [3] S. Rachev [et al.] , Financial models with Lévy Processes and Volatility Clustering , John Wiley & Sons, Inc., Hoboken, New Jersey 2011.
- [4] S. Raible, Lévy processes in finance: theory, numerics, and empirical facts (Ph.D. dissertation), Freiburg University, Germany, 2000.
- [5] Y. Sahalia, J. Cacho-Diaz, R. Laeven, Modeling financial contagion using mutually exciting jump processes, J. Bank. Finance. Econ. 117 (3) (2015) 586–606.
- [6] D. Hainaut, Clustered Lévy Processes and their financial applications, Journal of Computational and Applied Mathematics, 2017.

Fractional Stochastic Volatility Models and Applications in Pricing

Samira Amirian¹

Amirkabir University, Tehran, Iran

Erfan Salavati

Amirkabir University, Tehran, Iran

Abstract

We study the fractional stochastic volatility model in which the volatility is driven by a fractional Brownian motion and the price is driven by an independent simple Brownian motion. We relate the option price to a quadratic average of the exponential fractional Brownian motion and we derive the asymptotics of the mentioned average as t tends to infinity.

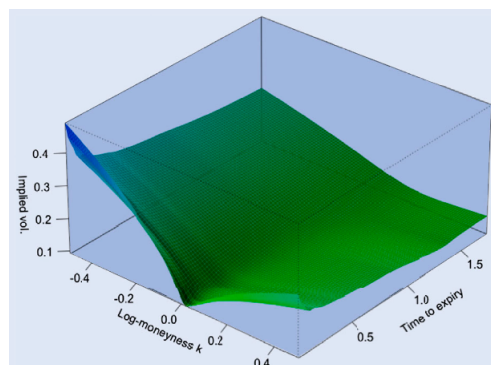
Keywords: Fractional stochastic volatility model, volatility smile, call option pricing, asymptotics of the distribution density

AMS Mathematical Subject Classification [2018]: 91G20, 60G22

1 Introduction

Implied volatility surface is the plot of the implied volatility (σ) as a two variable function of moneyness (strike price K) and expiration time (T). This surface is obtained from empirical data of option prices traded in the options markets. Figure 1 shows the volatility surface.

Figure 1: Implied volatility surface



¹speaker

Should Black-Scholes model be the ground truth of the market, the volatility surface would be flat (which is not the case). Hence the curvature of this surface is an indicator of how much the Black-Scholes model fits to the market.

So far, one of the challenges of mathematical finance has been to build more sophisticated models that illustrate the same implied volatility as observed in real data.

One of the efforts in this direction has been made using the so called stochastic volatility models (SVM). These models assume that the volatility is not constant but instead is a stochastic process in itself. Hence $\sigma(t)$ also follows a stochastic differential equation alongside the price process $S(t)$. One of the most famous such models is the following due to Hull and White (1987).

$$\frac{dS(t)}{S(t)} = \mu(t, S(t))dt + \sigma(t)dw(t)^1 \quad (1)$$

$$d(\ln \sigma(t)) = K(\theta - \ln \sigma(t))dt + \gamma dw(t)^2 \quad (2)$$

where w^1 and w^2 are independent Wiener processes. A simple argument shows that when conditioned on the trajectory of $\sigma(t)$, the price at time t of a European option of exercise date T is indeed the Black-Scholes price where the constant volatility σ is replaced by its quadratic average over the period $\sigma_t^2 = \frac{1}{T-t} \int_t^T \sigma^2(u)du$. Hence the option price can be obtained by taking expectation of this Black-Scholes price.

Although the SVM models fit better to the volatility surface, they are still far from a good fit. In recent years a new family of models have been introduced which is a generalization of SVM models in the sense that they use the fractional Brownian motion as the noise in the volatility process.

Definition 1.1. A fractional Brownian motion (fBm) with Hurst parameter $0 \leq H \leq 1$ is a zero-mean Gaussian process $(W_t^H)_{t \in R}$ with the covariance

$$E[W_t^H W_s^H] = \frac{\sigma_H^2}{2} (|t|^{2H} + |s|^{2H} - |t-s|^{2H})$$

The fractional stochastic volatility model (FSVM) is given by the following system:

$$\frac{dS(t)}{S(t)} = rdt + \sqrt{V_t}dB(t) \quad (3)$$

$$d(V_t) = \mu V_t dt + \sigma V_t dW_t^H \quad (4)$$

The same argument as in the SVM models implies:

Theorem 1.2 (see [2]). *The European call option price is given by*

$$C(t) = S(t)E_t^Q \left[\Phi \left(\frac{m_t}{U_t} + \frac{U_t}{2} \right) | F_t \right] - e^{-m_t} E_t^Q \left[\Phi \left(\frac{m_t}{U_t} - \frac{U_t}{2} \right) | F_t \right]$$

where

$$m_t = \ln \left(\frac{e^{-r(T-t)} S(t)}{K} \right), \quad U_t = \sqrt{\int_0^t \sigma^2(u) du}$$

and Φ is the standard Gaussian distribution function.

2 Main results

Our ultimate goal is to study the properties of the option prices in the fractional Hull-White model of the previous section. Theorem 1.2 shows that any information on the distribution of the variable $U(t)$ could be employed to obtain information on the distribution of the option price as well. Hence one can study the distributional properties of $U(t)$.

Article [3] does the same thing in the classical Hull White model and then uses it to study the asymptotic behaviour of the distribution density of the stock price process. Following the framework of [3] we define

$$\alpha_t = \int_0^t e^{\mu s + \sigma W_s^H} ds$$

And notice that by the time reversing property of the fBm we have,

$$\alpha_t \stackrel{d}{=} e^{\mu t + \sigma W_t^H} \int_0^t e^{\mu(s-t) + \sigma(W_s^H - W_t^H)} ds$$

Now we want to use the Ito formula. The Ito formula for fBm exists only for $H > \frac{1}{2}$ (except $H = \frac{1}{2}$ which is indeed the Bm itself). For $H > \frac{1}{2}$ the Ito formula is simply the chain rule. By applying Ito formula we find,

$$dV_t = (\mu V_t + 1) dt + \sigma V_t dW_t^H$$

let $f(v, t)$ be the density function of V . By using the fractional version of the Kolmogorovs forward equation we the following partial differential equation governing f :

$$\partial_t f + \frac{\partial}{\partial v}((2v + 1)f) - H t^{2H-1} \frac{\partial^2}{\partial v^2}(\sigma v f) = 0$$

$$f(t, 0) = f(t, \infty) = 0, \quad f(0, x) = \delta_{x_0}$$

Questions of interest, regarding the function $f(t, x)$ are its asymptotic behaviour when $t \rightarrow 0, \infty$ and also $x \rightarrow 0, \infty$.

In this article we provide an asymptotic bound for f when $t \rightarrow \infty$ and show that under certain assumptions, it decays exponentially and obtain the rate.

Theorem 2.1. *If for some $t_0 > 0$ and a positive constant M , we have $|f(t_0, x)| < M$, then*

$$f(t, x) \leq M e^{-\mu(t-t_0)}$$

For proof we use the result of [4] on asymptotic of the solution of hyperbolic PDEs.

References

- [1] Christian Bayer, Peter Friz and Jim Gatheral, pricing under rough volatility, Quantitative Finance, (2015).
- [2] FABIENNE COMTE, ERIC RENAULT, long Memory in continuous time Stochastic volatility Models, Mathematical Finance, Thesis, Imperial College, (1998).

- [3] Archil Gulisashvili, Elias M. Stein, Asymptotic behavior of the distribution of the stock price in models with stochastic volatility:the Hull and White model, Elsevier, (2006).
- [4] Chris Cosner, Asymptotic Behavior of solutions of second order parabolic partial differential equations with unbounded coefficients, Journal of Differential Equations, (1980).

e-mail: `samiraamiriy@aut.ac.ir`

e-mail: `erfan.salavati@aut.ac.ir`

S&P 500 Stock Selection via Factor Analysis and Principal Component Analysis

Alireza Fallahi¹, Erfan Salavati, Adel Mohammadpour

Department of Mathematics and Computer Science, Amirkabir University of Technology, Tehran, Iran

Abstract

The sheer size of data in the new age is not only a challenge for computer hardware but also the essential bottleneck for the efficiency of many machine learning algorithms. In this issue, we, describe the difference between factor models and principal component analysis (PCA). Both techniques are utilized to “simplify” complex data sets mainly collected of multiple time series as a purpose of a smaller number of time series.

Keywords: Principal Component Analysis, Factor Analysis, Stock Selection.

Mathematics Subject Classification [2018]: 62P05, 62H25

1 Introduction

Factor models and PCA have been widely employed in many applications, such as portfolio management (Fama and French, 1992; Carhart, 1997; Fama and French, 2015; Hou et al., 2015), large-scale multiple testing (Leek and Storey, 2008; Fan et al., 2012), high-dimensional covariance matrix estimation (Fan et al., 2008, 2013), and forecasting using many predictors (Stock and Watson, 2002a, b), high-dimensional covariance matrix estimation (Fan et al., 2008, 2013), corporate finance, performance management, and many other areas of financial analytics. In spite of this fact that factor models and PCA overlap many resemblances with linear regression analysis, there are also considerable differences. With factor analysis, we attempt to know that if and how we can describe our variables as a multiple linear regression on a reduced number of independent variables. We will begin with a definition of factor models and factor analysis and then introduce some ways of creating factor models. PCA is a historical method for mathematicians and is still one of the most common techniques for data analytics and visualization. High dimensional data, for example, images, macroeconomics dataset, often have the attribute whose it lies in a low-dimensional area and that many dimensions are highly correlated.

2 Assumptions

Categories of Factor Analysis

We can write explicitly factor model as form as N equations:

$$\begin{aligned} y_{1t} &= a_1 + b_{11}f_{1t} + \cdots + b_{1q}f_{qt} + \varepsilon_{1t} \\ &\vdots \\ y_{it} &= a_i + b_{i1}f_{1t} + \cdots + b_{iq}f_{qt} + \varepsilon_{it} \quad i = 1, \dots, N, \quad t = 1, \dots, T \\ &\vdots \\ y_{Nt} &= a_N + b_{N1}f_{1t} + \cdots + b_{Nq}f_{qt} + \varepsilon_{Nt} \end{aligned}$$

¹speaker

The a_i are the constant terms, the b_{jt} coefficients are called factor loadings, the f_{jt} are the hidden factors, and the ε_{it} are the error terms.

The main two assumption of factor models as following:

1. Zero-mean variables of factors and residuals
2. Assumption of uncorrelatedness between factors and residuals vectors, that is:

$$E(f_t) = 0, \quad E(\varepsilon_t) = 0, \quad E(f_t \cdot \varepsilon_t) = 0, \quad \text{for any } t$$

Categories of Principal Component Analysis

With algebra glance, principal components are specific linear combinations of the p random variables $X_1, \dots, X_p$. And also, in the Geometric approach, these linear combinations show the selection of a new coordinator or code selection that gains from $X_1, \dots, X_p$ as a code axis. The updated axes give us the highest volatility directions and also provide us with a more conservative structure of covariance. As we will see, principal components belong only on the covariance matrix (or correlation matrix) of $X_1, \dots, X_p$. Let the random vector $X' = [X_1, \dots, X_p]$ have the covariance matrix Σ with weights $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_p \geq 0$ let's see the relation of them:

$$\begin{aligned} Y_1 &= a'_1 X = a_{11}X_1 + a_{12}X_2 + \dots + a_{1p}X_p \\ Y_2 &= a'_2 X = a_{21}X_1 + a_{22}X_2 + \dots + a_{2p}X_p \\ &\vdots \\ Y_p &= a'_p X = a_{p1}X_1 + a_{p2}X_2 + \dots + a_{pp}X_p \end{aligned}$$

And then if we using this equation :

The linear combinations $Z = CX$ have

$$\begin{aligned} \mu_Z &= E(Z) = E(CX) = C\mu_X \\ \Sigma_Z &= \text{Cov}(Z) = \text{Cov}(CX) = C\Sigma_X C' \end{aligned}$$

Then we can have this result:

$$\begin{aligned} \text{Var}(Y_i) &= \mathbf{a}'_i \Sigma \mathbf{a}_i \quad i = 1, 2, \dots, p \\ \text{Cov}(Y_k, Y_i) &= \mathbf{a}'_i \Sigma \mathbf{a}_k \quad i, k = 1, 2, \dots, p \end{aligned} \tag{1}$$

$y_1, \dots, y_p$ that are the uncorrelated linear combinations and each one has the maximum variances as significant as it can. See equation (1).

The first definition has the greatest difference from the others that can be extracted from this formula:

$$\text{maximizes } \text{Var}(y_1) = \mathbf{a}'_1 \Sigma \mathbf{a}_1$$

First principal component = linear combination $\mathbf{a}'_1 \mathbf{X}$ that maximizes

$$\text{Var}(\mathbf{a}'_1 \mathbf{X}) \quad \text{subject to} \quad \mathbf{a}'_1 \mathbf{a}_1 = 1$$

Second principal component = linear combination $\mathbf{a}'_2 \mathbf{X}$ that maximizes

$$\text{Var}(\mathbf{a}'_2 \mathbf{X}) \quad \text{subject to} \quad \mathbf{a}'_2 \mathbf{a}_2 = 1 \text{ and}$$

$$\text{Cov}(\mathbf{a}'_1 \mathbf{X}, \mathbf{a}'_2 \mathbf{X}) = 0$$

At the i th step,

i th principal component = linear combination $\mathbf{a}'_i \mathbf{X}$ that maximizes

$$\text{Var}(\mathbf{a}'_i \mathbf{X}) \quad \text{subject to} \quad \mathbf{a}'_i \mathbf{a}_i = 1 \text{ and}$$

$$\text{Cov}(\mathbf{a}'_i \mathbf{X}, \mathbf{a}'_k \mathbf{X}) = 0, \quad \text{for } k < i$$

3 Application to Financial Data

Explanation of Data: The data are daily returns of $m = 470$ equities on the S&P 500 index from 04.01.2010 through 30.12.2016, for a total of 1762 observations. In this empirical example these time series should be transformed by taking logarithms and/or differencing to make them approximately stationary.

The scree plot method shows us a line for each factor and its eigenvalues. Number of the eigenvalues that are greater than one regarded as the number of factors.

In our dataset, the first "15" factors have the eigenvalue of more than 1 in the Factor Analyse test. It means that we necessity to choose only 15 factors (or unobserved variables). We can also calculate the factor loading (B) and the variance of the residuals (Ψ) over the assumptions that factors are uncorrelated.

See our 5 factors detail:

Table 1: The "SS loadings" row is the sum of squared loadings

	1	2	3	4	5
SS Loadings	327.582940	55.278911	25.944567	13.682766	10.637179
Proportion Variance	0.696985	0.117615	0.055201	0.029112	0.022632
Cumulative Variance	0.696985	0.814600	0.869801	0.898913	0.921545

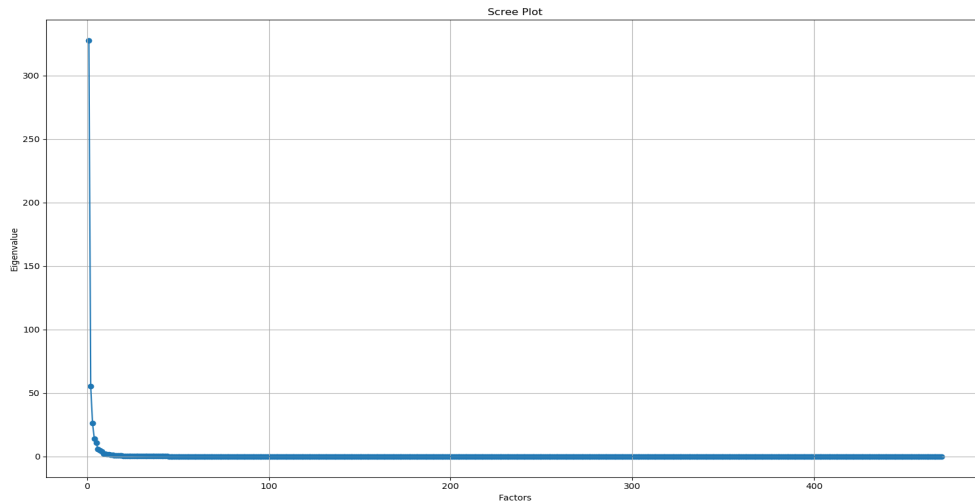
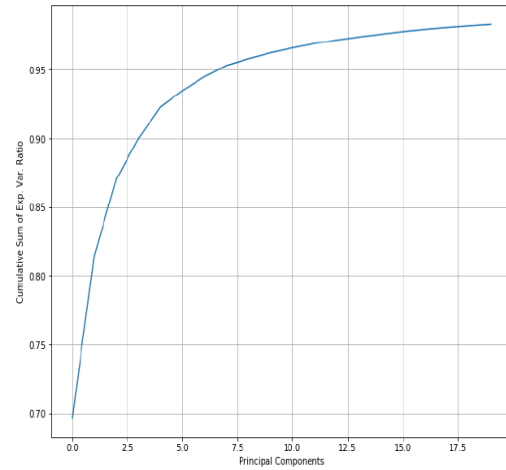
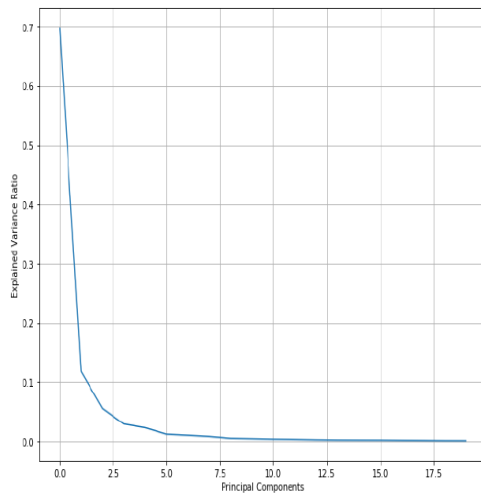


Figure 1: the first "15" factors have the eigenvalue of more than 1

Study of factor investigates the broad dataset and considers connections interlinked. It reduces the observed variables to a few unknown variables or defines the classes of interrelated variables that help market researchers compact market situations and identify the secret relationship between customer interest, desire, and cultural influence.

See our result from the PCA method:

The first "32" principal components cover 0.99 of the variance.



References

- [1] Rechar A. Johnson and Wichern. Dean W. *Applied Multivariate Statistical Analysis* , Pearson, Sixth Edition.
- [2] Fabozzi. Frank , *The Basics of Financial Econometrics* , Wiley.

Email: alirezafallahi@aut.ac.ir

Email: erfan.salavati@aut.ac.ir

Email: adel@aut.ac.ir

Non-life insurance reserve for solvency purposes

Fatemeh Atatalab¹
Shahid Beheshti University, Tehran, Iran

Amir Teimour Payandeh Najafabadi
Shahid Beheshti University, Tehran, Iran

Abstract

This paper considers the problem of predicting the claim reserve based on aggregated data from run-off triangles. We analyze claims development results and quantify its prediction uncertainty. This is an important view in solvency considerations and risk- based controlling of non-life insurers.

Keywords: Claim Development Result, Reserve, Mean Square Error of Prediction.

Mathematics Subject Classification [2018]: G22, C13

1 Introduction

Claim reserve is an estimate of an insurer's liability from future claims. Chain ladder is the simplest method to predicting claims reserves but this method only establishes a point estimate of claim reserve for accounting purposes which is insufficient for the actuarial analyst to understand the possible fluctuations in the claim reserve and their impacts on the entity's profit and loss statement, as well as on the balance sheet. In this paper we predict claim reserve by using four methods and compares them with the methods that provided by central insurance of Iran. The Solvency II Directive (2009/138/EC) is a Directive in European Union law that codifies and harmonises the EU insurance regulation. Primarily this concerns the amount of capital that EU insurance companies must hold to reduce the risk of insolvency. Solvency II is a far-reaching program of prudential regulations, which vary in severity depending on the riskiness and diversity of an insurer's business. Solvency capital requirement (SCR) is a Value-at-Risk (VaR) measure based on a 99.5% confidence interval of the variation over one year of the amount of "basic own funds" (broadly assets minus technical provisions). The SCR includes non-life insurance underwriting risks. One of the uncertainty sources in non-life insurance is estimating companies' liability especially claim reserve. In most solvency considerations one is interested into the changes and uncertainties over a one-year time horizon. That is, one predicts the outstanding loss liabilities today and in one year with the new information available in one year. The difference between these two successive predictions is the so-called claims development result (CDR). In international regulation framework, we calculate one-year CDR and conditional mean square error of prediction (MSEP). This is an important view in solvency II considerations.

1-1 Models and formulas

In this section, we introduce models for claim reserving that are based on triangle models. Also, we represent stochastic approach for calculating CDR and its prediction uncertainty.

1-1-1 Chain ladder

A way to classification of data is use run-off triangle. In this triangle, let noted accident year by i and development year by j . development year means number of years that claims have delay to reported or paid. The main task of actuaries is to predict the lower triangle. We assume that $C_{i,j}$ is cumulative claims that occur in year i and developed j year after, $0 \leq i \leq I$ and $0 \leq j \leq J$.

¹ speaker

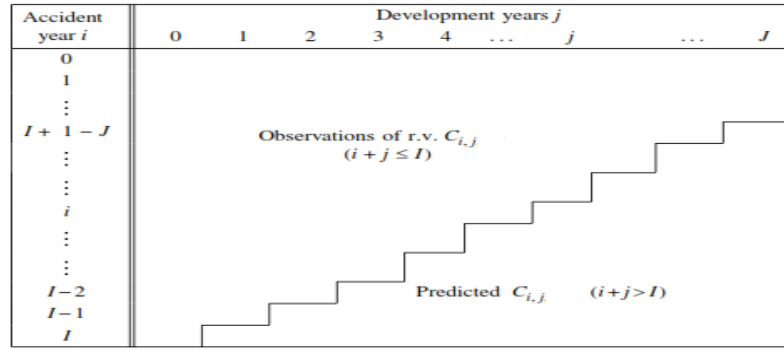


Figure1: run-off triangle

Model assumption 1 (Classical chain ladder):

- Cumulative claims $C_{i,j}$ of different accident years I are independent.
- There exist development factors $f_0, \dots, f_{J-1} > 0$ such that for all $0 \leq i \leq I$ and $0 \leq j \leq J$ we have

$$E[C_{i,j} | C_{i,0}, \dots, C_{i,j-1}] = E[C_{i,j} | C_{i,j-1}] = f_{j-1} C_{i,j-1}$$

Development factor is

$$\hat{f}_j = \frac{\sum_{i=0}^{I-j-1} C_{i,j+1}}{\sum_{i=0}^{I-j-1} C_{i,j}}$$

□

Claim reserve amount for accident year $i = 1, \dots, I$ and development year $j = I - i + 1$ is equal to $\hat{R}_i = \hat{C}_{i,j} - C_{i,j}$. Chain ladder algorithm isn't a stochastic model; hence we can't calculate uncertainty such as MSEP. To solve this problem stochastic models in chain ladder algorithm are developed, for example Mack distribution-free chain ladder (Mack, 1993), over-dispersion Poisson by England and Verrall (2002) and Bayesian chain ladder by Gisler and Wuthrich (2008). Mack distribution-free chain ladder is more usual between these methods because it is simple and distribution-free.

Model assumption 2 (Mack distribution-free chain ladder):

- Cumulative claims $C_{i,j}$ of different accident years I are independent.
- $(C_{i,j})_{j \geq 0}$ is a Markov process and exist $f_j > 0$ and $\sigma_j > 0$ for all $0 \leq i \leq I$ and $0 \leq j \leq J$ that

$$E(C_{i,j} | C_{i,j-1}) = f_{j-1} C_{i,j-1} \quad \text{and} \quad \text{Var}(C_{i,j} | C_{i,j-1}) = \sigma_{j-1}^2 C_{i,j-1}$$

□

1-1-2 Mean Square Error of Prediction (MSEP):

In this subsection we want to study the accuracy of predictions by model assumption 2. In claim reserving, general method to measure uncertainty is conditional MSEP. Assume that we are in time I and information D_I is available. We would like to estimate variability of reserves that is solvency II requirements. Wuthrich and Merz (2008) showed that MSEP under model assumption 2 is given by

$$mse_{p_{C_{i,j}|D_I}}(\hat{C}_{i,j}^I) = \text{Var}(C_{i,j} | D_I) + (E(C_{i,j} | D_I) - \hat{C}_{i,j}^I)^2 \quad (1)$$

1-1-3 Claim Development Reserve (CDR):

The study of the total uncertainty of the full run-off is a long-term view. Another important view is short-term view that is important for management decisions. Based on model assumption 2 and observation D_I in time I and D_{I+1} for one year after we have

$$E(C_{i,I} | D_I) = C_{i,I-I} \prod_{j=I-I}^{I-1} f_j \quad \text{and} \quad E(C_{i,I} | D_{I+1}) = C_{i,I-I+1} \prod_{j=I-I+1}^{I-1} f_j$$

For claims reserving at time I we predict the total ultimate claim with the information available at time I and, similarly, at time $I+1$ we predict the same total ultimate claim with the (updated) information available at time $I+1$ ($D_I \subset D_{I+1}$). The CDR at time $I+1$ for accounting year $(I, I+1]$ is then defined to be the difference between these two successive predictions for the total ultimate claim.

Definition 1: Based on model assumption 2, CDR for accident year I in accounting year $(I, I+1]$ is given by

$$CDR_i(I+1) = E(R_i^I | D_I) - (X_{i,I-I+1} + E(R_i^{I+1} | D_{I+1})) = E(C_{i,j} | D_I) - E(C_{i,j} | D_{I+1})$$

That $X_{i,I-I+1} = C_{i,I-I+1} - C_{i,I-I}$ is incremental of paid losses.

□

The conditional MSEP is then defined by

$$mse_{p_{CDR_i(I+1)|D_I}}(0) = \text{Var}(CDR_i(I+1) | D_I) = E(C_{i,j} | D_I)^2 \frac{\sigma_{I-I}^2 / f_{I-I}^2}{C_{i,I-I}} \quad (2)$$

For calculation of CDR we use bellow equation:

$$C\tilde{D}R_i(I+1) = \hat{R}_i^{D_i} - (X_{i,I-i+1} + \hat{R}_i^{D_i+1}) = \hat{C}_{i,j}^I - \hat{C}_{i,j}^{I+1}$$

1-1-4 Minimum-variance method for claim reserving:

In this method, for a given insurance portfolio, let denote the burning cost by M_{ih} on underwriting year i as observed at the end of h -th year of run-off. Cumulative of M_{ih} denoted by C_{ih} . Our observation in run-off triangle denoted by $D = \{C_{ih}; i=1, 2, \dots, n, h=1, 2, \dots, i\}$. Claims reserve for occurrence year q is:

$$E(E(C_{qm}|D) - \hat{\mu}_{qm})^2) = \text{minimum}$$

And $E(\hat{\mu}_{qm}) = E(C_{qm})$. We assume that after m years all claims are settled. $\hat{\mu}_{qm}$ is unbiased estimator for $E(C_{qm}|D)$.

Model assumption 3 (Minimum-variance method for claim reserving):

- 1) C_{ih} and $C_{i'h'}$ are stochastically independent for $i \neq i'$
- 2) $E(C_{ih}) = e_h$ independent of i .
- 3) $P_i \text{Cov}(C_{ih}, C_{i'h'}) = \sigma_{hh'}$ independent of i (P_i is premium volume of year i)
- 4) $\hat{\mu}_{qm} = \sum_{i=1}^n \sum_{h=1}^i \alpha_{ih} C_{ih}$

For more details see (Kramreiter and Straub, 1973). Claims reserve by using this method is

$$\hat{R}_q = \hat{\mu}_{qm} - C_{qq}$$

Theorem 1: based on model assumption 3, conditional MSEP is given by

$$MSEP_{C_{qm}|D}(\hat{\mu}_{qm}) = \text{Var}(C_{qm}|D)$$

Note that $\hat{\mu}_{qm} = E(C_{qm}|D)$ is an unbiased estimator. □

1-1-5 Credible loss ratio reserving method:

In this section we introduce credible loss ratio method based on (Werner, 2009). Werner's formulas recalled hereafter, because they are used in our numerical example. The considered credible loss reserving method requires slightly less information. We suppose that there are I underwriting periods, for which one knows besides actuarial premiums V_i , $i = 1, \dots, I$, used as a measure of exposure, m_k is loss ratio. X_{ik} is the incremental of losses. Loss ratio is $m_k = \frac{\sum_{i=1}^{I-k+1} X_{ik}}{\sum_{i=1}^{I-k+1} V_i}$. The quantity $U_i^{BC} = V_i \sum_{k=1}^{I-i+1} m_k$ is nothing else than the loss ratio estimate of the total ultimate

claims required for the underwriting period i . The loss ratio payout defined by $P_i = \frac{\sum_{k=1}^{I-i+1} m_k}{\sum_{k=1}^I m_k}$. Individual total ultimate claims amount is given by $U_i^{ind} = \frac{C_{i,I-i+1}}{P_i}$. Individual loss ratio IBNR claims reserve, is defined by

$$R_i^{ind} = q_i U_i^{ind}; \quad i = 1, \dots, I$$

Where $q_i = 1 - P_i$ represents the proportion of the total ultimate claims, which is expected to be paid in the future for the underwriting period i . On the other side, collective loss ratio IBNR claims reserve is

$$R_i^{coll} = q_i U_i^{BC}; \quad i = 1, \dots, I$$

The individual claim reserve considers the latest accumulated claims amount to be fully credible predictive for future claims and ignores the prior burning cost estimate of the total ultimate claims, while the collective claims reserve ignores the current accumulated paid claims and relies fully on this prior estimate. Therefore it is natural to apply the credibility mixture to those reserves and use the credible loss ratio IBNR claims reserve estimate

$$R_i^c = Z_i R_i^{ind} + (1 - Z_i) R_i^{coll}$$

Z_i is the credibility weight associated to the individual loss ratio reserve. It interesting to reconsider two popular choices of the credibility weights proposed in the literature. Gunnar Benktander(1976) proposed the credibility weight $Z_i^{GB} = P_i^{GB}$. This leads to the Benktander loss ratio IBNR claims reserve. Walter Neuhaus (1992) corresponds to the credibility weight $Z_i^{WN} = \sum_{k=1}^{I-i+1} m_k$. It leads to the Neuhaus loss ratio IBNR claims reserve.

1-1-6 Claims reserve prediction based on central insurance of Iran directive:

In year 2018, insurance regulator of Iran based on article 58 of third party liability (TPL) of automobile act and article 10 of technical reserve law, approves a directive for TPL loss reserve adequacy. In this section we modeling this directive and compare its result by previous method as mentioned already. Based on this directive, insurance companies oblige to calculate their liabilities in TPL line of business by using chain ladder or loss ratio based on partial information methods for companies with more than 5 years experience. The proposed chain ladder method is the same as model assumption 1. We focus on Loss ratio based on partial information method. Assume that accident year denoted by $i=1,2,3,4,5$ and development year is $j=0,1,2,3,4$.in this method we assume that losses are normalized with $M_{ij} = \frac{X_{ij}}{V_i}$ where V_i is premium. X_{ij} is incremental loss that occur in year i and develop until year j .

First of all we calculate the mean of each column. After that we calculate cumulative mean $C_j = \sum_{k=0}^j \frac{1}{I-k} \sum_{i=1}^{I-k} \frac{X_{ik}}{V_i}$.

The percentage of unearned liabilities is given by $N_j = C_j - C_j = \sum_{k=j+1}^J \frac{1}{1-k} \sum_{l=1}^k \frac{X_{lk}}{v_l}$, Where $J=4$. Claims reserves amount is $\hat{R}_i = P_i \hat{N}_{(I-i)}$; $\hat{R} = \sum_{i=I-i+1}^I \hat{R}_i$

Model assumption 4 (loss ratio based on partial information):

- 1) $X_{ij} = Z_i Y_j$, where Z_i and Y_j are independent. Assume that $\sum_{j=1}^J Y_j = 1$
- 2) X_{ij} are independent
- 3) $Var(X_{ij}) = P_i^2 \sigma_j^2 + (\sigma_j^2 + \gamma_j^2) \sigma_i^2$, $E(X_{ij}) = P_i \gamma_j$

□

2 Main results

In this section we use TPL loss data of one Iranian insurance company with more than 5 years experience. Results of calculation of claims reserves is given in table 1.

Table1. Claim reserve calculation (million IRR.)

year	Model 1&2	Model 3	Credible loss ratio- Benktander	Credible loss ratio- Neuhaus	Model 4
1391	0	0	0	0	0
1392	275122	273020	268939	268344	265382
1393	1288087	1274523	127753	1276028	1253270
1394	6485426	6425132	6414043	6360220	6036588
1395	23251594	23127068	26438829	26647884	27551922
Total reserve	31300229	31099743	34399564	34552474	35107162

Table1 show that the reserve for models 1 and 2 are the same. Chain ladder method is suitable for number of loss prediction and provide weak estimate of loss amount. In credible loss ratio method, current cumulative loss paid and prior information consider together. In model 4 we can't predict lower triangle's cells. Central insurance proposed two methods for calculation of claim reserving by insurance companies with more than 5 experiences. As can be seen in table1, these two methods (1 and 4) have around 4 billion IRR differences.

Root of conditional uncertainty prediction of CDR based on model assumption 2 is 97% total uncertainty. The reason of this high value is that knowing the next diagonal (I+1) in claim development triangle already releases a major part of the claims run-off risks. The amount of prediction uncertainty for model assumption 3 is equal to 97% too. Claim development result present in table 2. The result show profit in insurer's balance sheet.

Table2. Claims development results (million IRR.)

Accident year	$E(C_{i,j} D_i)$	$E(C_{i,j} D_{i+1})$	CDR
1391	0	0	0
1392	275122	0	275122
1393	1288087	311880	976207
1394	6485426	1384794	5100632
1395	23251594	6459555	16792039
Total	31300229	8156229	23144000

With solvency II, insurers should not just predict the reserve, but also asses the uncertainty of their predictions. Next to solvency II, also the upcoming IFRS17 regulation will encourage insurers to get more detailed grasp on their reserve. We propose to Iranian insurance industry to use more detail for predicting their loss reserve and use international regulation recommendation to calculate loss reserve based on micro model and calculate the uncertainty of predicting and claims development result.

References

- [1] G. Benktander, *An approach to credibility in calculating IBNR for casualty excess reinsurance*, Actuarial Review, (1976), pp. 7.
- [2] T. Mack, *Distribution-free calculation of standard error of chain ladder reserve estimates*, ASTIN Bulletin, 23 (1993), pp. 213-225.
- [3] W. Neuhaus, *Another pragmatic loss reserving method or Bornhuetter/Ferguson revisited*, Scandinavian Actuarial Journal, (1992), pp.151-162.
- [4] H. Werner, *Credible loss ratio claims reserves: the Benktander, Neuhaus and mack method revisited*, ASTIN Bulletin, 39(2009), pp. 81-99.
- [5] M.V. Wuthrich, M. Merz, *stochastic claims reserving method in insurance*, Wiley finance series, 2008.

Email: f_atatalab@sbu.ac.ir

Email: amirtpayandeh@sbu.ac.ir

A Forecasting Model for Limit Order Book in Tehran Stock Exchange Using LSTM Network

Parisa Davar¹

Institute for Advanced Studies in Basic Sciences, Zanzan, Iran

Parvin Razzaghi

Institute for Advanced Studies in Basic Sciences, Zanzan, Iran

Ali Fourosh Bastani

Institute for Advanced Studies in Basic Sciences, Zanzan, Iran

Abstract

In this paper, we propose a deep learning based model first presented in [J. Sirignano, R. Cont, *Universal features of price formation in financial markets: perspectives from deep learning*, Quantitative Finance, 19 (2019), pp. 1449–1459], to predict price movements in Limit Order Book (LOB) data of the Tehran Stock Exchange (TSE). Our model is trained on data from 30 top companies, utilizing stacked Long Short-Term Memory (LSTM) to capture longer time dependencies. Importantly, experimental results show that the average accuracy of the model is 88.80% in out-of-sample data set. It demonstrates the model's ability in extracting universal features in the considered financial market.

Keywords: Deep Learning models, Time series forecasting, Limit Order Book data, Stock analysis

1 Introduction

Recently, machine learning models, especially deep learning models, have gained much popularity for a variety of tasks such as image analysis, natural language processing and speech recognition. Deep learning model is a powerful method for extracting features. Recurrent Neural Networks (RNNs) are a type of deep learning based models mostly used for forecasting time series data. RNNs have a challenge memorizing knowledge due to the vanishing gradient property and Long Short-Term Memory (LSTM) was introduced to solve this challenge. During the last years, these models have achieved remarkable results in financial modeling such as market forecasting [3], deep hedging [1] and risk management [2]. In this paper, we focus on forecasting stock market using limit order book data, because it is an important task for market practitioners and academia as an investment management tool.

Among the existing approaches, theoretical models provide a mathematical and economic understanding of the limit order book dynamics. However, in these models, there are some assumptions which limit it in

¹speaker

using in the real markets. In contrast, Sirignano [3], shows that deep neural networks achieve the better performance compared to other methods such as decision trees, boosted trees, random forests and support vector machines in data-driven models of the limit order book. More recently, Sirignano and Cont [4], pooled all data from different stocks (independent of the specific asset) to improve generalization ability of deep learning based model in price formation in financial market which is called the universal model.

In this paper, we present a universal deep learning based approach to model stock price movement in Tehran Stock Exchange. Also, we find out that the preprocessing step is necessary to get reliable performance. To this end, we preprocess the data with MinMax scaler normalization and smoothing method.

2 Main Problem

Our main goal is to estimate the price formation mechanism. It maps an order flow (other variables can be used) to the market price which is defined as:

$$\text{Price}(t + \Delta t) = F(\text{Price history}(0...t), \text{Order flow}(0...t), \text{Other information}), \quad (1)$$

$$= F(X_t, \epsilon_t), \quad (2)$$

where X_t is the state of the limit order book at time t and ϵ_t is a random noise.

2.1 Limit Order Book data structure

A limit order is an order to buy or sell a certain number of an asset at a specific price. Unexecuted limit orders are placed in Limit Order Book (LOB) based on the price and time levels. In this paper, we use five levels of both price and volume on each side (buy and ask) of the limit order book from the 30 top companies for 24th March 2012 to 9th November 2013 between 08:30:00 and 12:30:00 as training dataset. These 30 top companies are among the most liquid stocks listed from 2012-2013. So, we have 20 features at each timestamp and specifically, each data is a time series of 100 timestamps of LOB in the form:

$$X = [x_1, x_2, \dots, x_{100}]^T \in R^{100 \times 20}, \quad (3)$$

$$x_t = [p_a^i(t), p_b^i(t), v_a^i(t), v_b^i(t)]_{i=1}^{n=5}, \quad (4)$$

where X is a single input, p_a^i and p_b^i are respectively, an ask and a bid price and v_a^i and v_b^i are an ask and a bid volume at i th level of a limit order book. It should be noted that, we have normalized the LOB data for each asset by MinMax scaler of scikit-learn package.

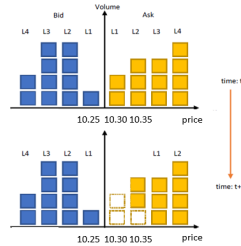


Figure 1: Limit Order Book example at time t and $t + 1$.

In a given financial market, mid-price is an average of the best bid and ask prices on the first level of LOB which represent the direction of price changes. So, we use mid-price movement as our labels. Because

of the high frequency data structure, we use smoothing labelling method of the form:

$$S(t) = \frac{1}{k} \sum_{i=0}^k P_{t-i}, \quad (5)$$

as a preprocessing step in which P_i is the i th mid-price. Then, we consider $M_t = S_{t+1} - S_t$ to provide the direction of price movement at time t . Therefore, if $M < 0$ or $M > 0$, we define it as down or up and if $M = 0$, there is no difference between the two consecutive smooth prices which it means as a stationary state.

2.2 Model Architecture

In this section, the model architecture is provided. Our model consists of three stacked LSTM layers and two dense layers with rectified linear units (ReLUs) and softmax as activation function on the top stacked LSTM layers, as shown in Figure (2). The last layer transforms the output to a probability distribution for the next price move. LSTM is a sequential model which has a memory cell and three gates as input gate, output gate and forget gate. They could maintain its state over time and store useful information to model long term features in sequential data. Stacked LSTM is a variation of LSTM which consists of more than one layer of LSTM. It learns a function from sequence of past observations to map it to an output by minimizing a loss function. In this paper, each LSTM layer has 100 hidden units and the output size of it is 20. We use the categorical crossentropy loss function and Adaptive Moment Estimation (ADAM) as the optimization algorithm. Also, we set the learning rate to 0.001.

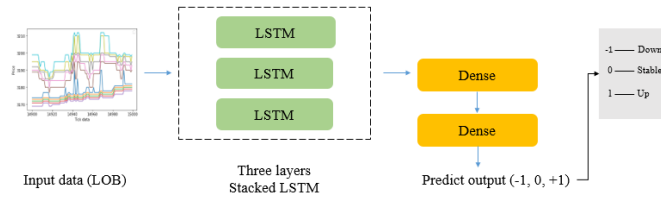


Figure 2: Model architecture schematic.

3 Main results

In this section, we illustrate the results of applying the proposed method on the 30 top Iranian companies. We split the data into two groups: train and test set. Training is done on limit order book for the 30 stocks from 24th March 2012 to 22th August 2013 with 20% of the data considered as the validation data. We assess the proposed model on the in-sample and out-of-sample data (10 companies from different industries) from 24th August 2013 to 9 November 2013. In this experiment, there are two versions of our approach. In the first version, it is trained on each specific stock and in the second model, it is trained on a pool of all stocks. The second model is referred to as the universal model. To evaluate the model, two performance measures are utilized. The first measure is an accuracy criteria and the second one is F1 which is a measure of test's accuracy. Table (1) presents the results of our model for out-of-sample data. We report the results for in-sample-data and stock-specific cases in table (2). Also, We display the confusion matrices in Figure (3). Our results demonstrate the model's ability to extract universal features and it confirms that there is a universal price formation mechanism from the dynamics of supply and demand in a financial market.

Table 1: Experimental results on transfer learning

Company	Accuracy %	F1 %	Company	Accuracy %	F1 %
Mokhaberat Iran	89.71	89.77	Machinesazi Niro Mohareke	84.67	84.68
Karton Iran	88.98	88.96	Palayesh Bandarabas	88.20	88.15
Keshtirani Iran	96.65	96.45	Sanaye Lastiki Sahand	86.36	86.34
Sanaye Fazar Ab	91.06	91.01	Group Mapna	91.53	91.50
Lole va Machinesazi Iran	86.33	86.30	Tose Sakhteman	84.57	84.34

Table 2: Experimental results for the in-sample data

Model	Average accuracy %
Universal	86.99
Stock specific	48.98

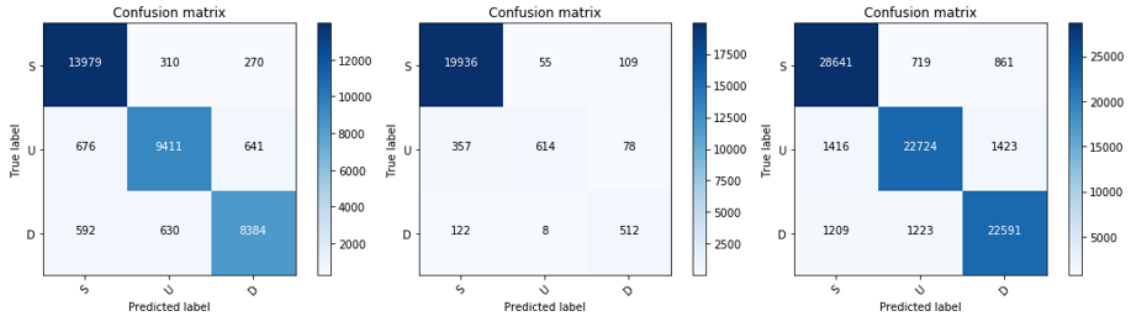


Figure 3: Confusion matrices. Results on Sanaye Fazar Ab, Keshtirani Iran, and Group Mapna from the left to right.

References

- [1] H. Buehler, L. Gonon, J. Teichmann and B. Wood *Deep hedging*, Quantitative Finance, 28 (2019), pp. 1–21.
- [2] A. Petropoulos, V. Siakoulis, E. Stavroulakis, A. Klamargias *A robust machine learning approach for credit risk analysis of large loan level datasets using deep learning and extreme gradient boosting*, IFC Bulletins Chapters, 49 (2019).
- [3] J. Sirignano, *Deep learning for limit order books*, Quantitative Finance, 19 (2019), pp. 549–570.
- [4] J. Sirignano, R. Cont, *Universal features of price formation in financial markets: perspectives from deep learning*, Quantitative Finance, 19 (2019), pp. 1449–1459
- [5] Z. Zhang, S. Zohren, S. Roberts, *DeepLOB: Deep convolutional neural networks for limit order books*, IEEE Transactions on Signal Processing, 67 (2019), pp. 3001–3012

e-mail: pdavar@iasbs.ac.ir
e-mail: p.razzaghi@iasbs.ac.ir
e-mail: bastani@iasbs.ac.ir

Pricing Catastrophe bonds under a jump-diffusion interest rate model by Finite element method

Pegah Amiri¹

Department of Mathematics, AllamehTabataba'i University, Tehran, Iran.

Abdolsadeh Neisy

Department of Mathematics, AllamehTabataba'i University, Tehran, Iran.

Abstract

A catastrophe bond (CAT) is a financial instrument designed to provide money for insurance companies in the occurrence of devastating natural disasters like earthquakes or tsunamis. If a predetermined catastrophe or event occurs, the investors will lose the principal they invested, and the issuer will receive that money to cover their losses. The aim of this paper is to model the price of the catastrophe risk bonds. under the assumption that the occurrence of the catastrophe is independent of financial market behavior, we employ a jump-diffusion interest rate model and solve the resulted PIDE with the finite element method. Also, we consider a dependency between the claims sizes and the claim inter-arrival times for the aggregate claims as a semi-Markov process and obtain explicit CAT bond prices formulae with regard to four different payoff functions.

Keywords: Pricing CAT bond, Jump-diffusion interest rate, Markov-dependent environment, Finite element method.

1 Introduction

When natural and man-made disasters occur, losses and recovery costs are typically covered by a combination of utility companies, special insurance programs and/or governments. For instance, mainly losses from the 2011 Fukushima disaster were covered by the government of Japan. Because, in such cases, financial demands on insurance and reinsurance businesses are potentially enormous, introducing a securitization method to achieve an adequate liquidity fund is significant. An alternative method is issuing Catastrophe (CAT) bonds. CAT bonds transfer the financial consequences of catastrophic events from the risk carrier (an insurance company, country or a regional government) to investors in a contract. Insurers and reinsurers typically issue cat bonds through a special purpose vehicle (SPV), a company set up specifically for this purpose. Cat bonds pay high-interest rates and diversify an investor's portfolio because natural disasters occur randomly and are not correlated with other economic risks.

In this study, we derive CAT bond pricing formulae under the assumption that the occurrence of a catastrophe is independent of the global financial market behavior. To obtain a complete model, we consider the CIR interest rate with a jump. The financial markets may receive information either through small, gradual perturbations or large, sudden shocks (Merton, 1976). For interest rates, jump-diffusion processes are particularly meaningful since the interest rate is an important economic variable, which is, to some extent, controlled by the government as an instrument. These jumps are caused by several market phenomena such as money market interventions by the Fed, news surprises, and shocks in the foreign exchange markets, and so on. In pricing financial derivatives, term structure models with jump are particularly important because ignoring jumps in financial prices may cause

¹ speaker

inaccurate pricing. Furthermore, similar to [1], we formulate our model in a Markov-dependent environment and model the dependency between the claim inter-arrival times and the claims sizes for the aggregate claims as a semi-Markov process. The advantage of this consideration is the development of a more realistic model, where both the claim severity and occurrence time before the next claim are partially dependent on the claim intensity, which indicates the seasonality effect of catastrophe events. For instance, major catastrophe event like an earthquake may trigger many other disasters (e.g., tsunami, flood, side earthquakes, wildland fire, and landslide) in a short time. Also, by considering four different payoff functions (classical zero-coupon and coupon, multi-threshold zero-coupon, and defaultable), we obtain analytical formulae for catastrophe bonds.

2 Mechanism of catastrophe bonds

The general structure of Catastrophe Bonds is presented in Figure 1. SPV enters into a reinsurance agreement with a sponsor or counterparty (e.g., insurer, reinsurer, or government) by issuing CAT bonds to investors and receives premiums from the counterparty or sponsor in return for providing a pre-specified coverage. This agreement specifies the conditions and terms for when the insurance will cover a claim and which triggers will activate the policy. The existence of SPV minimizes the frictional cost of capital and eliminates the counterparty risk. Therefore, sponsors can transfer part of the risks to investors who willingly accept the risk in exchange for higher expected returns. The SPV collects both the premiums and the principle (respectively received from sponsors and Investors) and invests in a collateral account, where they are typically invested in highly-rated money market funds. The returns generated from collateral accounts are swapped for floating returns based on LIBOR (London Interbank Offered Rate). The investors' coupon payments are made up of SPV investment returns and the premiums the sponsor pays. Before the maturity time of the CAT bond, if no trigger event occurs, the collateral is liquidated at the maturity date, and investors are paid compensation for bearing the catastrophe risks plus their principle (solid line in Figure 1). Whereas, if before the maturity a trigger event occurs, the SPV will

liquidate collateral to make the payment and reimburse the sponsor under the terms of the catastrophe bond agreement, and CAT bond investors will only receive part of their principle (dashed line in Figure 1).

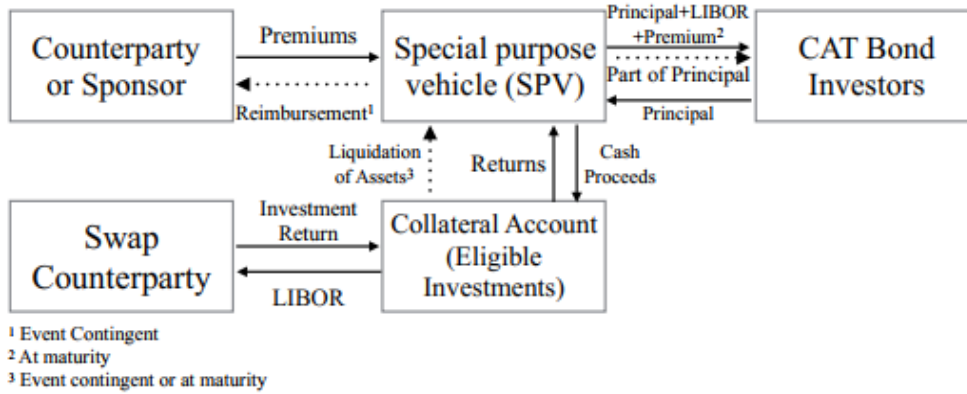


Figure 1

3 Modeling the catastrophe bond

In this paper, we use an insurance industry index trigger. The aggregate loss process is modeled as a compound distribution process defined by the frequency (inter-arrival times) and the severity (claim sizes) of disasters.

Classical Cramer–Lundberg risk model, stated that risk models are characterized by two stochastic processes: claim number process and claim amounts process. In this study, we assume these two processes are mutually independent. But in order to consider a more appropriate model, in modeling the aggregate losses, Similar to Shao (2015), we add dependence between the claim sizes and the inter-arrival times in the claims process.

Furthermore, we develop an PIDE representation for r-component of the proposed CAT bond model. we consider the CIR interest rate with a jump and the short-rate dynamics $\{r(t): t \in [0, T]\}$ can be expressed as follows:

$$dr = \mu(r)dt + \sigma(r)dw + Jdq \quad (1)$$

$$\mu(r) = k(\theta - r(t)), \quad \sigma(r) = \sqrt{r(t)}$$

with the condition $2k > \sigma^2$ where $\theta, k, r(0)$, and σ are positive constants. Jump size J is a normal variable with mean μ and standard deviation γ

When interest rates follow the SDE(1), a bond has a price of the form $B(r, t, T)$. We set up a riskless portfolio and the jump-diffusion version of Ito's lemma to functions of r and t . And then, we derive the partial differential bond pricing equation.

From Ito's Lemma, we can write dB as

$$dB = \{B_t + B_r \mu(r) + \frac{1}{2} \sigma^2 B_{rr}\} dt + \sigma(r) B_r d\omega + \{B(r+J, t) - B(r, t)\} dq \quad (2)$$

By considering two non-identical CAT bonds B_1 and B_2 , We can construct a risk free portfolio by hedging B_1 with B_2 . So we obtain the following PIDE

$$B_t + [\mu(r) - \xi_r \sigma(r)] B_r + \frac{1}{2} \sigma^2(r) B_{rr} - (r + \lambda) B + \lambda M = 0 \quad (3)$$

Where $M = E^Q[B(r+J, t)]$.

We use finite element method to solve this PIDE.

We consider four types of payoff functions (the zero-coupon, the multi-threshold zero-coupon, the defaultable zero-coupon, and the coupon payoff functions) for CAT bonds with T maturity time. Under the assumption that the payoff function is independent of the financial risks variable, by means of these payoff functions, the structure of interest rate and aggregate loss for a predetermined threshold level D , we obtain four models for the value of a catastrophe bond at time t using the standard tool of a risk-neutral valuation measure.

4 Solution

We employ finite element method for solving the following PIDE for r -component of the CAT bond model,

$$B_t + [\mu(r) - \xi_r \sigma(r)] B_r + \frac{1}{2} \sigma^2(r) B_{rr} - (r + \lambda) B + \lambda M = 0$$

$$M = E^Q[B(r+J, t)]$$

The basic idea of the finite element method is approximating the solution of a differential equation with a set of algebraically simple functions by dividing the spatial domain of the differential equation into sub-domains called elements. The parameters of this function for each element are usually different. The functions are equal regarding the function type yet different concerning the values of the parameters. Each of these functions only has local support. To put it into a more vivid picture, outside a small number of elements, it takes on the value zero. The elements are no overlapping and cover the domain on which the differential equation is defined.

FE can be used not only to find approximate solutions for a given differential equation but also any equation of calculus, hence integral, integro-differential, and variational equations can be solved as well. J. Topper (2005).

5 Main results

In this section, we price CAT bonds using the standard tool of a risk-neutral valuation measure with four different payoff functions for T time maturity one-period CAT bonds.

In an arbitrage-free market, at time t , the price of an attainable contingent claim with payoff $\{P(T) : T > t\}$ can be shown by the fundamental theorem of asset pricing. By assuming the payoff function is independent of the financial risks variable (interest rate):

$$V(t) = E^Q \left(e^{-\int_t^T r(s) ds} P_{CAT}(T) | F_t \right) E^Q (P_{CAT}(T) | \mathcal{F}_t) = B(t, T) E^P (P_{CAT}(T) | \mathcal{F}_t).$$

1) Let $V^{(1)}(t)$ be the price of the T-maturity zero-coupon CAT bond under the risk-neutral measure Q at time t with payoff function $P_{CAT}^{(1)}(t)$ for a predetermined threshold level D, the value of CAT bond at time t is given by:

$$V^{(1)}(t) = B_{CIR}(t, T) z(\eta + (1 - \eta) F(T - t, D)),$$

where $F(T - t, D)$ represents the accumulated function of the aggregate loss.

2) Let $V^{(2)}(t)$ be the price of the T-maturity zero-coupon CAT bond under the risk-neutral measure Q at time t with payoff function $P_{CAT}^{(2)}(t)$ (with a multi-threshold value) for a predetermined threshold level D, the value of CAT bond at time t is given by:

$$V^{(2)}(t) = B_{CIR}(t, T) Z \sum_{k=1}^h \eta_k (F(T - t, D_k) - F(T - t, D_{k-1}))$$

3) Let $V^{(3)}(t)$ be the price of the T-maturity coupon CAT bond under the risk-neutral measure Q at time t with the payoff function $P_{(3)}^{CAT}$ (with a coupon payment at the maturity date) for a predetermined threshold level D, the value of CAT bond at time t is given by:

$$V^{(3)}(t) = B_{CIR}(t, T) (Z + CF(T - t, D))$$

4) let $V^{(4)}(t)$ be the price of the T-maturity coupon CAT bond under the risk-neutral measure Q at time t with the payoff function $P_{(4)}^{CAT}$ (defaultable payoff function) for a predetermined threshold level D, the value of CAT bond at time t is given by:

$$V^{(4)}(t) = B(t, T) Z \left[\tau + (1 - \tau - \tilde{F}(Z) - \tau \tilde{F}(\tau Z)) F_l(T - t, D) \right] + P \tilde{F}(\tau Z)$$

where $\tilde{F}(x)$ denotes the issuing company's default probability at time T.

In this paper at first, we presented some concepts and mechanisms of Catastrophe Bonds. We developed a contingent claim process to price CAT bonds using models with a risk-free spot interest rate under assumptions of a no-arbitrage market, independently of the financial risks and catastrophe risks. With the concept of Financial Mathematics i.e., Ito formula and free risk portfolio, we modeled these bonds. By using the finite element method, we solved the resulted PIDE for r-component. Then, Under the risk-neutral pricing measure, bond price formulae are derived for four types of payoff functions when the trigger is determined by the aggregate loss process with a semi-Markov-dependent structure.

References

- [1] M Safaei, A Neisy, N Nematollahi, New Splitting Scheme for Pricing American Options Under the Heston Model, Computational Economics 52 (2018)405-420
- [2] J. Shao, A. D. Papaioannou, A. Pantelous, Pricing and simulating catastrophe risk bonds in a Markov-dependent environment, Applied Mathematics and Computation 309 (2017) 68–84
- [3] J. Topper, Financial Engineering with Finite Elements (2005)

Email: pegah_amiri@atu.ac.ir

Email: a_neisy@atu.ac.ir

Insured's Risk prediction in life insurance: Saman insurance case study

Reza Ofoghi

ECO College of Insurance, Allameh Tabataba'i University, Tehran, Iran

Fatemeh Ahangar Saryazdi

ECO College of Insurance, Allameh Tabataba'i University, Tehran, Iran

Abstract

Risk assessment is a crucial element in the life insurance business to classify the applicants. Companies perform underwriting process in order to make decisions on applications and to price policies accordingly. Increase in the amount of data and advances in data analytics, the underwriting process can be automated for faster processing of applications. This study aims at providing solutions to enhance risk assessment among life insurance companies using predictive analytics. Machine learning tools aim to gain this solution. The dataset consists of insured's purchase Saman life insurance from 2008 to 2017, with 7 attributes, which describe characteristics of life insurance insured's. Decision tree and multinomial regression use to predict insured's risk level. Multinomial regression shows that employment type is not significant in the variables. Decision tree plot also indicate employment type is not significant.

Keywords: Risk assessment, Life insurance, Predictive analytics, Decision tree, Multinomial Regression.

1 Introduction

The role of insurance in providing financial protection in the economy is well established. Over the years, the provision of cover by insurance companies has been crucial to the consummation of business plans and, by extension, wealth creation.

Risk assessment leads the insurance market to recognize the risk factors that might be occur insured's loss and suggest an instruction to assess new buyers as simplification of selling life insurance. It is also beneficiary to assess risk profile of every customer and offer the appropriate life insurance policy according to their risk level.

There are multiple types of non-financial risks: Hazard risk, operational risk and strategic risk. Strategic risk is closely related to the firm's overall strategies. Reputation risk, competition risk and regulatory risk are included in the strategic risk.

Data mining can be defined as the process of selecting, exploring and modeling large amounts of data to uncover previously unknown patterns. In the insurance industry, data mining can help companies gain business advantage. By using data mining techniques insurance companies can extract their customer's financial behavior and also predict and design appropriate policy to attract significant portion of market.

The findings of this research will redound to the benefit of insurance company considering that predictive analytics plays an important role in insurance risk assessment. The greater purchasing life insurance at the same time with customer profiling information rise accuracy and construct a guidance in order to control insured's risk.

This study is an attempt to elaborate on and predict insured's risk that purchase life insurance policy with considering insured's sex and age and their related policy features such as critical sum-insured, employment type.

In this case study we use insured's data set from 2008 to 2017 in order to answer these questions: (1) What factors effect insured's risk? (2) What is the each attributes coefficient value in risk level relationship? (3) What is the insured's risk prediction of life insurance policy?

Definition 1.1. Suppose that we wish to classify an observation into one of K classes, where $K \geq 2$. In other words, the qualitative response variable Y can take on K possible distinct and unordered values. Let π_k represent the overall or prior probability that a randomly chosen observation comes from the k th class; this is the probability that a given observation is associated with the k th category of the response variable Y . Let

$f_k(x) \equiv P_r(X = x | Y = k)$ denote the density function of X an observation that comes from the k th class.

Then Bayes' theorem states that

$$P_r(Y = k | X = x) = \frac{\pi_k f_k(x)}{\sum_{l=1}^K \pi_l f_l(x)}$$

Definition 1.2 A classification tree is classification tree used to predict a qualitative response. A classification tree, predict that each observation belongs to the most commonly occurring class of training observations in the region to which it belongs.

Since we plan to assign an observation in a given region to the most commonly occurring class of training observations in that region, the classification error rate is simply the fraction of the training observations in that region that do not belong to the most common class:

$$E = 1 - \max_k (\hat{p}_{mk})$$

Here $\hat{p}_{mk}$ represents the proportion of training observations in the m th region that are from the k th class. In a node m , representing a region R_m with N_m observations,

$$p_{mk} = \frac{1}{N_m} \sum_{i \in R_m} I(y_i = k)$$

Example 1.3. The dataset consists of insured's purchase Saman life insurance from 2008 to 2017, with 7 attributes, which describe characteristics of life insurance insured's. The data set comprises of nominal, continuous, as well as discrete variables. This study is an attempt to elaborate on and predict insured's risk that purchase life insurance policy with considering insured's sex and age and their related policy features such as critical sum-insured, employment type.

In this case study use insured's data set to answer these questions: (1) Which attributes effect insured's risk? (2) What is the each attributes coefficient value in risk level relationship? (3) What is the insured's risk prediction of life insurance policy?

Table1: Attributes Description

Attribute	Type	Description
Payment type	Categorical	3 payment type to pay premium
Employment-Info	Categorical	3 Employment to category insured's risk
Policy period	numeric	Normalized insured's policy period
Insured's sex	Categorical	2 types of Insured's sex
Insured's age	numeric	Normalized insured's age
Sum insured	numeric	Normalized Sum-insured price
Critical illness sum insured	numeric	Normalized Critical illness sum-insured price
Response	Categorical	6 categorical insured's risk level

Solution. Decision tree and multinomial regression are the model use to answer example 1.3.

Main results

The following table is a decision tree result model.

Confusion matrix is the evaluation matrix in decision tree model which is indicate in following table.

	Predicted 1	2	3	4	5	6	Error
Actual							
1	9698	2709	221	7	0	166	24.2
2	5122	4606	156	3	0	104	53.9
3	1021	315	377	9	0	50	78.7
4	1238	402	69	26	0	95	98.6
5	802	80	169	7	0	58	100
6	949	175	40	4	0	335	77.7

Overall error: 48.1%, Averaged class error: 72.18333%

Figure 1 shows that insured's second risk level portion is 44 percent and third risk level is 34 percent portion. In second risk level insured's age less than 46 years old have the 47 percent. In the six and seventh risk level insured's age greater than 46 years old have the most portion among other insured's.

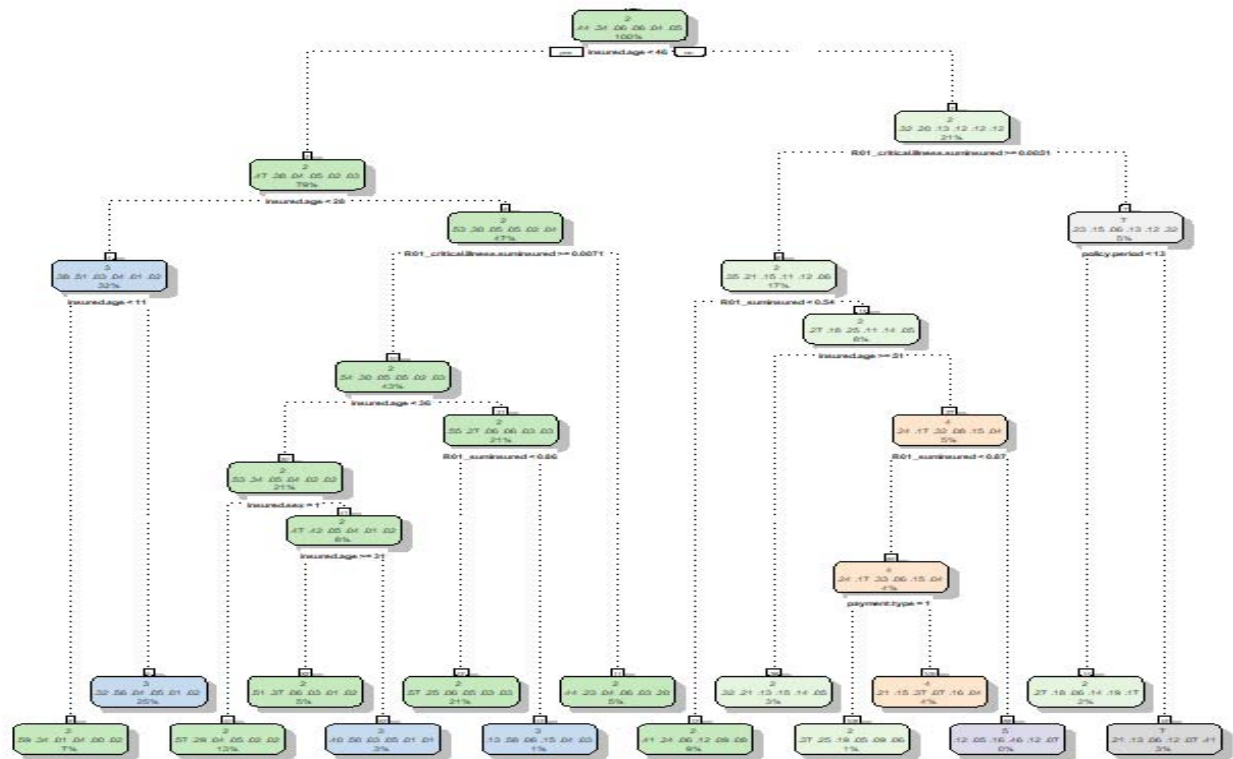


Figure 1: A Decision tree model

Another model applied to dataset is multinomial regression with the following result.

variable	LR Chisq	Df	Pr(>Chisq)
Policy period	81.56	5	3.955e-16 ***
Employment type	15.53	10	0.114
Insured age	1822.07	5	< 2.2e-16 ***
Insured sex	216.49	5	< 2.2e-16 ***
suminsured	173.39	5	< 2.2e-16 ***
Critical illness suminsured	719.97	5	< 2.2e-16 ***

As the result in Table3 all of the variables except employment type are significant. The evaluation AIC is AIC: 73210.09 in this model, and Pseudo R-Square: 0.06394640. Descriptive analysis also applied to this data set.

Acknowledgment

We thank Mr. Parviz Moradi, Saman insurance actuary manager Dr. Shima Ara and Saman actuary experts who provided insight and expertise that greatly assisted the research, although they may not agree with all of the interpretations/conclusions of this paper.

References

- [1] Noorhannah Boodhun¹, Manoj Jayabala, *Risk prediction in life insurance industry using supervised learning algorithms*, (2018),.
- [2] Trevor Hastie, Robert Tibshirani, Jerome Friedman, *The Element of Statistical*, University of London, 2nd ed., Springer, 2017.
- [3] Raja R V Kulkarni, *Applications of Data Mining Techniques in Life Insurance*, 2012.
- [4] S. Balaji, S. K. Srivatsa, *Naïve Bayes Classification Approach for Mining Life Insurance Databases for Effective Prediction of Customer Preferences over Life Insurance Products*, 2012.

Email: ofoghi@atu.ac.ir

Email: fatemeh.ahangar.s@alumni.ut.ac.ir

Fraud detection methods in health insurance Case study: An Iranian insurance company

Oreinab Afrooz Keldarshi¹
Razi Insurance Company, Tehran, Iran

Saman Ebrahimpour
Razi Insurance Company, Tehran, Iran

Abstract

As we know, Insurance fraud is an important and costly problem for both policyholders and insurance companies in all sectors of the insurance types such as health insurance. Fraud in health insurance is done by intentional deception or misrepresentation the facts for gaining more benefits and it is vary due to inadequacy of relevant laws or influence of community culture. In this paper, we introduce type of fraud in medical expenses in several countries, then we investigate frauds by health care providers. Data mining tools and techniques can be used to detect fraud in large sets of insurance claims data. These tools are divided into two learning techniques including supervised and unsupervised that is employed to detect fraudulent claims. Specifically, in this paper we are going to use k-means algorithm for preventing fraud in health insurance industry.

Keywords: Insurance fraud, Anomaly detection, Data mining, Supervised and Unsupervised learning.

1. Introduction

Every year the expenditure healthcare is being exceeded by many of the countries. Due to the extreme growth of market size and their influential factors this application domain requires a high-end data analytics mechanism. The significant problem of this wing is fraud, waste, abuse includes improper billing, repeated claims, uncovered services, drug abuses, counterfeit drugs, off-label marketing issues and many more. There is a series of technical challenges for data analytics. As there is massive storage of data over a period of time and from the representation point of view these all are many diverse datasets [P.Naga Jyothi, et al \(2019\)](#).

Fraud is one of the major problems which cause significant losses in insurance industry. Insurance fraud can be referred to "Inflating Loss". [Gill et al. \(1994\)](#) defined fraud in the insurance industry as "knowingly making a fictitious claim, inflating a claim or adding extra items to a claim, or being in any way dishonest with the intention of gaining more than legitimate entitlement." Although the amount of losses incurred through fraud in insurance industry is difficult to analyze, such losses are more common in contracts with higher premiums. Researchers have shown that investigating insurance fraud is often difficult and not cost-effective, because failing to detect fraud incorrectly can hurt honest customers and delay their payments. High costs involved in detecting fraud are also worrisome, so many insurers pay for claims without proper investigation.

¹ speaker

Sometimes the laws that apply to fraud in some countries are not specific to insurance fraud and do not adequately reduce fraud in the industry. Lack of specific laws makes it too difficult to prosecute criminals. With growth of social health and medical insurance, insurance industry has become one of the biggest victims of frauds, and this may accelerate development of a comprehensive legal framework to control and supervise fraud in health insurance. Therefore, with increasing rates of frauds in many countries, traditional methods have not been effective in addressing this widespread problems, and therefore, data mining tools and techniques can be used to detect fraud in large sets of insurance claims data. In this paper, we are going to introduce other studies on detecting frauds in health insurance and employ machine learning methods to investigate frauds in health insurance data of an Iranian insurance company. (Melih Kirlidog, 2012).

Fraudulent health insurance is simply defined as: "A fictitious claim by which an insurer or health care provider is forced to contract or pay fictitious damages. This deliberate act results in financial gain for insured or medical center and is brought under false pretenses."

In the definition provided by the Health Insurance Association of America (HIAA), fraud and insurance fraud are deliberate misrepresentation, deliberate deception, or present of evidence by an individual or entity with knowing that this would result in unauthorized benefits and frauds. Examples of common abuses in health insurance are: additional tests to diagnose illness, prolonged hospital stay longer than needed, patient admission at night instead of day, patient admission only for symptoms and diagnosis, etc. So at first phase we identified and investigated some losses that were suspect to frauds. In what follows, we use unsupervised methods such as K-means to detect frauds. For this purpose we used health losses in an insurance company of Iran. We worked on dentistry losses for one years and found some losses that were suspect to frauds.

2. Data

In this paper we have worked on health insurance data during 2 years 2019-2020. We selected dental losses that were paid for two-years period because the most paid losses related to dental losses. There are different variables gathered from losses and the history of insured including national insurance code, health center, loss amount, loss announcement date, loss payment date, contract end date and the number of payments for each national code. The top 10 loss in health centers are shown in the figure below.

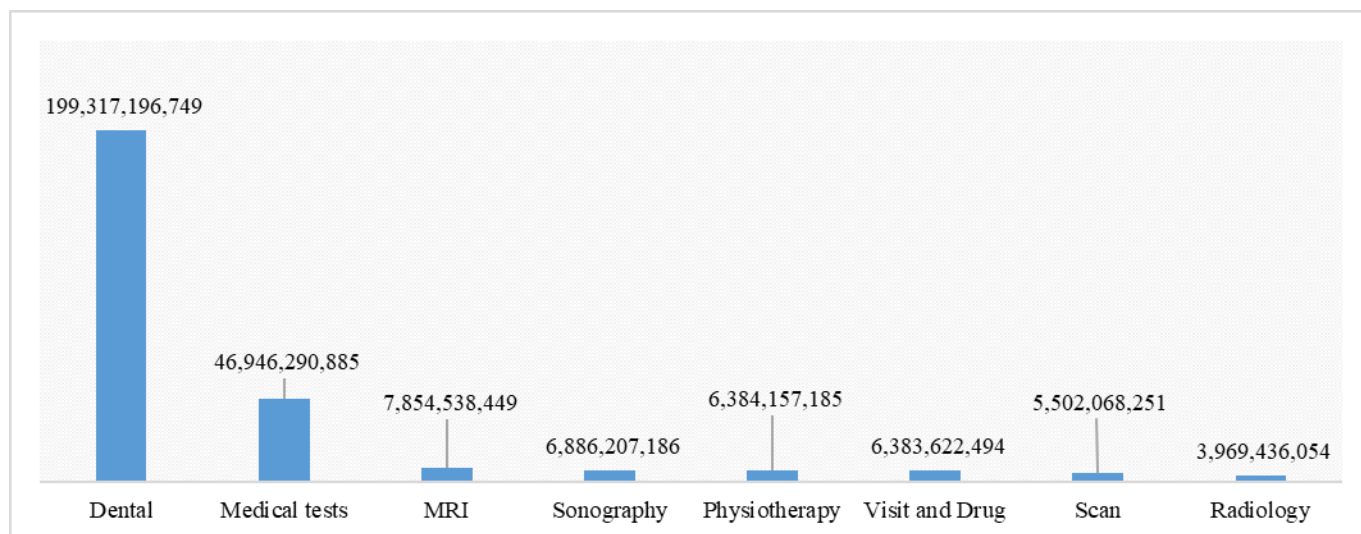


Figure 1: The Top 10 loss in health centers

The number of dental losses that were paid for two years are about 25,000 which are belong to 19,000 insureds. All insureds have been referred to the contract health centers.

Table 1: Summary statistics for claims data

Number of losses	Total losses amount (Rials)	Number of health centers
25,059	199,317,196,749	359

3. Analyzing methodology

There are two learning approaches in data mining models. Those are supervised method and unsupervised method: Supervised learning is the machine learning task of learning a function that maps an input to an output based on example input-output pairs. It infers a function from labeled training data consisting of a set of training examples.

We don't use this method in this paper, because we don't have labeled fraud data records in our country.

In this paper we use unsupervised method. Our goal is to find fraud patterns in unlabeled data. It focuses on finding those instances which show unusual behavior. Indeed we are going to discover both old and new types of fraud since they are not restricted to fraud patterns which already have pre-defined class labels like supervised learning techniques do.

Since fraud is high in health insurance especially in health centers, our focus is on losses that paid by health centers. Anomaly detection analysis was performed using K-means method. K-means clustering algorithms aims at partitioning n observations into a fixed number of k clusters. The algorithm will find homogeneous clusters.

Due to the random initialization, one can obtain different clustering results. When k-means is run multiple times, the best outcome, i.e. the one that generates smallest total within cluster sum of squares (SS), is selected. The total within SS is calculated as:

For each cluster results:

- for each observation, determine the squared related to distance from observation to center of cluster
- sum all distances

In this study we consider 10 clusters that are based on the amount of losses and national code, the figure is shown below.

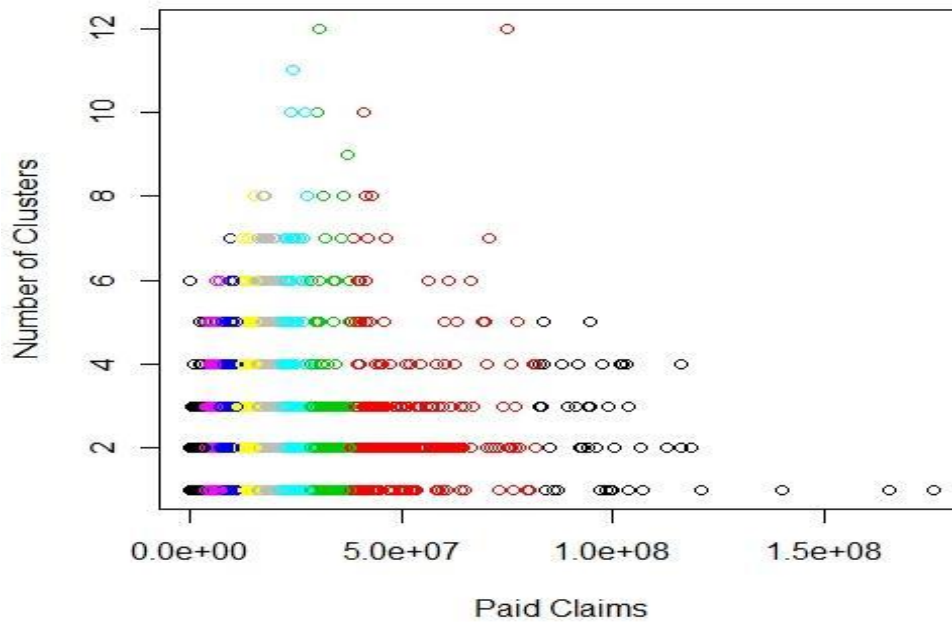


Figure 2: Paid losses clusters using K-means

For determine the number of clusters use these three steps

- Run k-means with $k=1, k=2 \dots k=n$; k is the pre-defined number of clusters.
- Record total within SS for each value of k .
- Choose k at the elbow position, as illustrated below. (Ramasubramanian A. S., 2017)

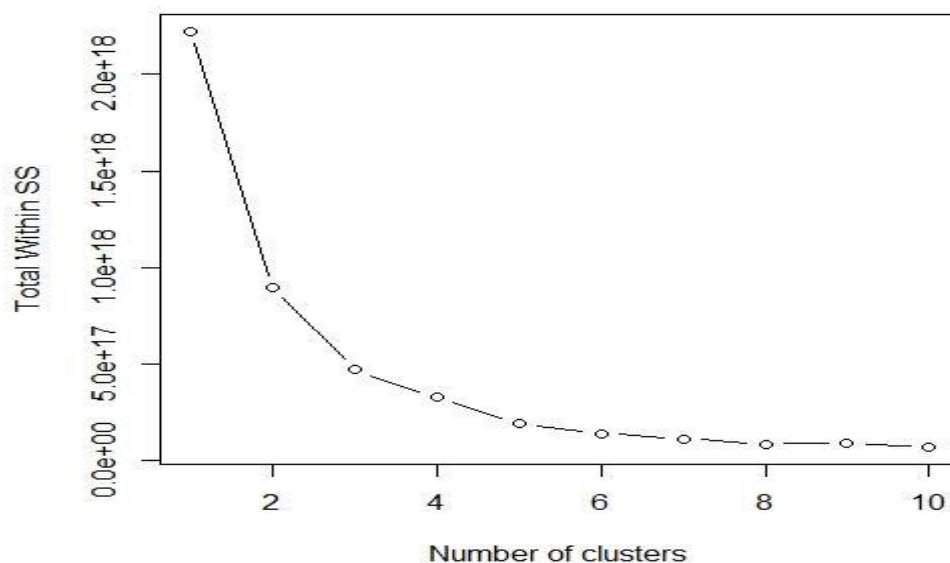


Figure 3: Total within square (SS)

In this study, losses above 60 million Rials were separated for each national code. We calculate mean difference between loss announcement date and loss payment date. In 13 health centers where have much more losses than others the average difference between loss announcement date and loss payment date was less than 1 day, the average payment per person in these centers is over 80 million Rials. We identified 66 losses suspected to fraud. These losses were paid by 13 health centers.

Table 2: Health center, average claims amount (Rials), and claim numbers

Health center name	Number of losses (for each national code)	Average of paid loss	Sum of paid loss	Mean difference between date of announcement and date of payment in days
Health center 1	11	76,867,727	845,545,000	0
Health center 2	9	88,895,444	800,059,000	0.3
Health center 3	9	88,282,833	794,545,500	1.1
Health center 4	6	78,300,667	469,804,000	0.6
Health center 5	5	70,000,660	350,003,300	1.0
Health center 6	4	78,364,400	313,457,600	0
Health center 7	4	66,249,250	264,997,000	0.1
Health center 8	3	90,345,667	271,037,000	0.0
Health center 9	3	84,411,333	253,234,000	1.8
Health center 10	3	83,139,333	249,418,000	0.8
Health center 11	3	84,689,667	254,069,000	0
Health center 12	3	105,461,667	316,385,000	0
Health center 13	3	74,471,667	223,415,000	1.0

4. Conclusion

Data mining techniques such as anomaly detection and clustering, can successfully detect anomalies or frauds in large datasets. This can be very useful for the insurance industry that is struggling with fraudulent claims. After identifying anomalous claims, further analysis should be carried out to find frauds. Fraud patterns are often recognized by insurance professionals, such research can reveal some new and unknown patterns.

In this article we find 13 health centers with maximum paid losses that showed inappropriate patterns in the payment process using data mining and clustering techniques. In these health centers, the average time between announcement

and payment date was less than one day. Note that these health centers are suspected of fraud and further investigation should be done to ensure frauds.

In conclusion, this paper reviews various methods for finding fraudulent behavior in health insurance claim. By analyzing the learning techniques, we will get a clear idea for the future work in health insurance claim fraud detection.

References

- [1] Melih Kirlidog, C. A. A fraud detection approach with data mining in health insurance, *Social and Behavioral Sciences* 62 (2012) 989 – 994.
- [2] Dallas Thornton, G. v. Outlier-based Health Insurance Fraud Detection for U.S. Medicaid Data, (2014).
- [3] Hossein Joudaki, A. R.-B. Improving Fraud and Abuse Detection in General Physician Claims: A Data Mining Study. (2016).
- [4] P.Naga Jyothi, D Rajya lakshmi, K.V.S.N.Rama Rao. Performance on Fraud Detection in Medical Claims of Healthcare Data. (2019).
- [5] Pravin R. Bagde, M. S. Analysis of Fraud Detection Mechanism in Health Insurance Using Statistical Data Mining Techniques, (2016).
- [6] Kathiresan, V and Gunasekaran Dr. S. and Faseela V. S, Health Insurance Claim Fraud Detection: A Survey, (2015).
- [7] Reichmann M, W. C. Individual and Institutional Corruption in European and US Healthcare: Overview and Link of Various Corruption Typologies, (2018).
- [8] Ramasubramanian K. and Singh A. Machine Learning Using R, (2017).

Email: Oafruz@gmail.com

Email: Saman_ebr@hotmail.com

Pricing Insurance Bonds Using RBF Method

M. Karimnejad Esfahani¹

Allameh Tabataba'i University, Tehran, Iran

A. Neisy

Allameh Tabataba'i University, Tehran, Iran

S. De Marchi

University of Padova, Padova, Italy

Abstract

In this paper, we consider the problem of pricing Catastrophe bonds which are based on interest rate and aggregate loss. Since both of these underlying assets are stochastic variables, so we need to select an appropriate and reasonable process for them. Concerning interest rate, different models with the feature of mean reversion have been proposed before, Cox-Ingersoll-Ross or Vasicek model, to name but a few. However, considering a suitable process for aggregate loss is not as clear as the first. It's partly due to the fact that the aggregate loss will not have continuous value necessarily. To put it simply, after occurring a catastrophe, the aggregate loss make some jumps from an specified value to another. Therefore, in order to bring this feature into account, it does make sense to assume the aggregate loss follows a Geometric Brownian Motion process with a jump-diffusion term. In spite of these explanations, the main focus of this work is on extracting a pricing model and the numerical solution for that. The latter is of great importance, mainly because, after going through constructing the desired model, the existence of jump-diffusion term leads us to a Partial-Integro Differential Equation (PIDE) for which no closed form solution can be proposed. On the other hand, dealing with Cat-Bonds, it's likely to consider additional underlying assets which means the dimension of the PIDE goes high. So, it seems reasonable to look for a method which is dimension-blind and because of this, in this research Radial Basis Function (RBF) is proposed as the numerical method. In addition, unlike other common method such as Finite Difference or Finite Element, no grid of points are required. In simple term, scatter points are enough for interpolating the solution. In fact, aforementioned features make this method attractive not only for solving an equation but also in other fields such as, machine learning and neural networks.

Keywords: Cat-Bonds, Partial-integro Differential Equation, Radial Basis Functions, Meshless Methods, Neural Networks

AMS Mathematical Subject Classification [2018]: 13D45, 39B42

1 Introduction

Generally, Cat-Bonds are based on at least 2 underlying assets: The aggregate loss L and the interest rate r . Very simply, the mechanism of Cat-Bonds include a sponsor who is looking for transferring and hedging

¹speaker

the risk, an investor who is looking for higher rate of return with the knowledge of higher risk, and a special purpose vehicle (SPV) which is a link between the first and the latter. It means that SPV provide the sponsor and investors with coverage againsts possible catastrophes and a chance to gain higher rate of return, respectively, by issuing Cat-Bonds. To this extend, the sponsor pays a premium to the SPV, so, in case a natural catastrophe happen with certain magnitude, in specified place and time, the SPV cover the damage which is inflicted to the sponsor. On the other hand, an investor purchase the bonds in exchange of a certain price with hope to gain more money at the expiry time T . After all the capital from both sponsor and investors come to the SPV, this vehicle invest all this fund and usually swapped the return with LIBOR rate using swap securities. Now, if nothing happens, the principle of the wealth, plus a repayment goes for the investor. It should be noted that, the latter is the reward of accepting risk of occurring a catastrophe. Moreover, if a catastrophe of great magnitude which is mentioned in the contract, happens then, the SPV pays the sponsor according to the contracts. Again, the reader should note that, the catastrophe should occur in the period and location that is clearly mentioned in the contracts. However, the structure of some Cat-Bonds are more complicated and pay some coupon to the investor. In addition, it is worth bearing in mind that, the way of determining the exact amount of delivered damage to the sponsor which is the result of catastrophe, plays a key role in these contracts. The reason is simply that, triggering a Cat-Bonds depends on the magnitude of inflicted loss to the sponsor.

In this work, with regards to the preceding part, we assume the price of Cat-Bonds V is a function of time, t , L and r . such that we have:

$$dL = \alpha dt + \sigma_L dW + \eta dq^Q \quad (1)$$

Here Q shows the martingale measure, η is the amount of jump that happen. Though, several model can be choosen for the interest rate, here for sake of simplicity, we assume it follows a simple GBM model. Thus:

$$dr = \mu dt + \sigma_r dW \quad (2)$$

Now, with respect to the Ito's lemma and the basic assumption that any value of bonds at any time can be traded in the market, the desired model will be achieved of the form:

$$V_t + (\alpha - \sigma_L)LV_L + \frac{1}{2}\sigma_L^2 L^2 V_{LL} + (\mu - \sigma_r)rV_r + \frac{1}{2}\sigma_r^2 r^2 V_{rr} - (r + \lambda^Q)V + \lambda^Q J \quad (3)$$

such that: λ is the market price of risk, λ^Q is the probability of occurring a jump and J is:

$$J = E[V(L + \eta, r, t)] \quad (4)$$

which is the integral part of the model.

In the preceding arguments we have mentioned RBF method as a desired numerical method for solving equation [3]. So here, we bring its definition and will discuss how it work solving a PIDE.

Definition 1.1. A function $\Phi : R^s \rightarrow R$ is called radial provided there exist a univariate function $\phi : [0, \infty) \rightarrow R$ such that:

$$\Phi(x) = \varphi(r) \quad r = \|x\| \quad (5)$$

and $\|\cdot\|$ is some norm on R^s .

In simple term, the definition 1.1 says that

$$\|x_1\| = \|x_2\| \Rightarrow \varphi(x_1) = \varphi(x_2) \quad (6)$$

One of the main advantages of this method in interpolation problems is that regardless of the dimension of the space, we use an univariate function instead of multivariate function and this feature is clearly due to the existence of the norm. Therefore, if one point of the space is a vector of the form $x = (x_1, x_2, \dots, x_s)$, with help of norm inside the radial function we would have $r = \|x\|$ which is a scalar.

Using this property, for interpolating function P at the point x we rewrite the function as follows:

$$P(x) = \sum_{k=1}^N a_k(t) \varphi(\|x - x_k\|_2) \quad (7)$$

such that x_j are the scatter points in our space. As we mentioned earlier, there is no need for a mesh or a grid of data and only scatter points are enough for interpolating a function. Because of this fact we also call this method meshless methods.

Now, if we put equation [7] in [3] the PDE part will easily change to a first order, ordinary differential equation of time t . However, the existence of the answer will be under question due to the fact that, there is no guarantee the resulted matrix would be invertible matrix. we can overcome this difficulty if the radial function in equation [7] would be positive definite. Thus, definitely, our linear system will have an answer.

2 Main results

In this research, an effort has been made to present a model for pricing Cat-Bonds based on two underlying assets while we assume that one of them, called aggregate loss, is exposed to some jumps which is the direct result of occurring a catastrophe. After extracting the pricing model, RBF method was briefly introduced as a proper numerical method for our work. One of the main reason for using RBF is that, in case we consider more underlying assets than we did in this paper, the dimension of the question goes high which means it would be a problem to deal with that because, most common numerical methods like Finite Difference or Finite Element will not properly work in high dimensions problems. On the contrary, as it's been mentioned, RBF method will not face difficulties in these kind of problems at least from theoretical point of view. Eventually, All these discussion means that, RBF method could be an excellent instrument leading us to more and more accurate pricing model of Cat-Bonds.

References

- [1] A. Neisy, M. Peymani, *Financial Modeling by Ordinary and Stochastic Differential Equations*, World Appl Sci J., 13 (2011), pp. 2288–2295.
- [2] S. De Marchi, M. Mandara, A. Viero, *Meshfree Approximation For Multi-Asset European and American Option Problems*, 1st edition, (2012)
- [3] Andrea J.A Unger, *Pricing Index-Based Catastrophe Bonds Part 1: Formulation and Discretization Issues Using a Numerical PDE Approach*, Computers Geosciences., 36(2) (2010), pp 139-149,.

e-mail: mohammad7394@gmail.com

e-mail: a_neisy@atu.ac.ir

e-mail: demarchi@math.unipad.it

The Relationship between Behavioral Orientation and Investor Behavior Based on Earnings Forecast Changes with Behavioral Finance Approach

Ebrahim Kargar

PhD in Economics, CEO of Dana Insurance

Mehdi Ahrari

PhD student in Economics, Allameh Tabatabaie University, Director of Dana Insurance Research and Development Bureau

Farzaneh Anvari¹

PhD in Economics Student of Tehran University, Expert of Dana Insurance Research and Development Bureau

Abstract

This paper investigates the effect of the behavioral bias on investors' financial decision making (insurance companies) in Tehran Stock Exchange for the period 2011-2017. For this purpose, panel data methods in the form of the single-factor model of the capital asset (including risk premium) and Fama-French three-factor model (including the size and ratio of book-to-market (B/M) as well as risk premium), and Carhart four-factor model (including momentum as well as risk premium, size and ratio of book-to-market (B/M)) were used. The results showed that the inversion of the return in the companies with consistently good performance is greater than the companies with inconsistently good performance. Also, the results showed that the reaction of the return in case of not being confirmed by the market in the companies with consistently good performance is greater than the companies with inconsistently good performance.

Keywords: changes in return prediction, behavior orientation, financial decision, panel data method.

1 Introduction

The course of the history of financial literature shows that at the beginning of the global economy growth in the early decades of the 20th century followed by the appearance and development of the money and financial markets, the pundits and scholars in the field of financial theories typically tried to explain the details of investment in the financial markets in the framework of the technical, economic and financial variables.

The traditional financial perspective assumes that people make rational decisions to maximize their wealth at a definite level of risk and to minimize the risk at a definite level of wealth. Such an approach stating the way people should behave is called normative. This perspective has provided the tools required to develop the portfolio theory, capital asset pricing, Arbitrage Pricing Theory and Pricing Option Theory. As opposed to this traditional approach, the financing based on the behavioral approach, perceptual and emotional errors that often influence the financial decision-makers which leads them to make undesirable decisions. This approach emphasizes the positive description of human behavior, and it investigates and studies the way people behave in practice in a certain financial field.

¹speaker

The financial theories have adopted two different approaches over recent past decades. The first approach: this is the neoclassic approach in the financial sciences whose fundamental financial theories' assumption-consists of the market efficiency and rational behavior of the investors in the market. This approach in the financial sciences started along with the capital asset pricing model and efficient market theory proposed in 1960s, and the midterm capital asset pricing and Arbitrage Pricing Theory (APT) of Miller and Modigliani in 1970s. The researchers found out over time and through different research that many people do not behave rationally and the existing chaos in the market cannot be justified using efficient market theory. Many cognitive and behavioral factors play a role in the investors' decision-making process. This leads to a behavioral revolution in financial issues through the paper presented by Kahneman and Tversky in 1979.

2 Research methodology

At first, the information on the net profit (loss) and operating profit and loss, and the sales of the listed companies in the Tehran Stock Exchange over ten years is collected. Then, the growth of these variables is calculated and the companies are classified into five classes of the same weight according to the growth of each one of these ratios. Then, the companies at the two ends of the classifications are classified into two portfolios, that is, purchase portfolio and sales portfolio where the difference in the return of the two portfolios for the periods under study are calculated, and the hypotheses are tested using the one-factor capital asset pricing model (CAPM) and Fama–French three-factor model and Carhart four-factor model are tested based on the difference in the performance of the two portfolios. In this research, augmented Dicky Fuller test; Kolmogorov–Smirnov test, error term normality test, White test and Durbin–Watson test were used as well as the descriptive statistics like the mean, standard deviation, and skewness.

In the present research, the data were gathered using library method where the researcher conducts the study through studying the relevant books and papers and collecting the information from the financial statements of the listed companies in the Tehran Stock Exchange through the information software called Tadbir Pradaz-Rahavard Novin. The present research tries to investigate all companies included in the definition of the statistical population. Thus, this research uses a time sampling method that tries to generalize the results to the periods other than the period under study

3 Conclusion and suggestions

For this purpose, the regression method in the form of the single-factor model of the capital asset (including risk premium) and Fama–French three-factor model (including the size and ratio of book-to-market (B/M) as well as risk premium), and Carhart four-factor model (including momentum as well as risk premium, size and ratio of book-to-market (B/M)) were used. The results showed that after a long period of time (5 years) of good performance, the company's return is inversed due to representative bias. Finally, the difference in the portfolios of consistent companies and inconsistent companies is calculated. The results showed that the inversion of the return in the companies with consistently good performance was greater in comparison with the companies with inconsistently good performance. Finally, to test the third research hypothesis, at first the financial performance of the companies is categorized according to hypothesis 2, in the next stage, the growth rate of the future period of the companies is examined. In this stage, the return of the securities based on the growth rate in the future period that confirms or contradicts the consistency (inconsistency) trend in the past is noted. The results showed that the return reaction, in case the trend is not confirmed by the market, in the companies with consistently good performance is greater in comparison with inconsistently good performance.

References

- [1] . Badri, A. Shavakhi, A. (2010). Reference Points, Stock Price and Volume of Transactions: The Evidence from Tehran Stock Exchange, Quarterly Journal of Stocks Exchange Market (12th ed.). 14-35.

- [2] . Javanmardi, M. Pourmousa, A. (2015). The Effect of Companies' Reports on Daily Behavior of Tehran Stock Exchange Market, *Empirical Accounting Research* (4th ed., Vol.4). 85-105.
- [3] . Jahangirirad, M., Marfou, M. Salimi, M. (2014). Investigation into the Group Behavior of the Investors in Tehran Stock Exchange Market, *Experimental Studies on Financial Accounting*. (11th ed. Vol.42).4. Heydarpour, F., Tariverdi, Y. Mehrabi, M. (2013). The Effect of Investors' Emotional Behaviors on Stock Return. *Quarterly Scientific Journal of Financial Knowledge of Securities Analysis*. (17th ed., Vol. 6).
- [4] . Saeedi, A. Farhaniyan, S. (2011). *Economic Fundamentals and Financial Behavior*. Tehran: Stock Exchange Notification and Service Company. 6. Shahabuddin, Sh. EsfandiyariMoghaddam, A. (2017). *The Effect of Herding Behavior* Tehran: Stock Exchange Notification and Service Company
- [5] . Shahabuddin, Sh. EsfandiyariMoghaddam, A. (2017). The Effect of Herding Behavior on the Performance of the Investing Companies Based on the Modern and Postmodern Theories of Portfolio, *Financial Research*, (1st ed., Vol. 19). 118
- [6] Abarbanell, Jeffrey S., and Victor L. Bernard, 1992, Tests of analysts' overreaction/underreaction to earnings information as an explanation for anomalous stock price behavior, *Journal of Finance* 47, 1181-1207.
- [7] Ang, Andrew, Robert Hodrick, Yuhang Xing, and Xiaoyan Zhang, 2006, The cross-section of volatility and expected returns, *Journal of Finance* 61, 259-299.
- [8] Baker, Malcolm, Robin Greenwood and Jeffrey Wurgler, 2008, Catering through nominal share prices, Working Paper, Harvard Business School
- [9] 0 10. Barberis, Nicholas, Andrei Shleifer, and Robert Vishny, 1998, A model of investor sentiment, *Journal of Financial Economics* 49, 307-343.
- [10] 1 11. Bartov, Eli, Suresh Radhakrishnan, and Itzhak Krinsky, 2000, Investor sophistication and patterns in stock returns after earnings announcements, *Accounting Review* 75, 43-63.
- [11] 2 12. Bernard, Victor L., and Jacob K. Thomas, 1989, Post-earnings-announcement drift: Delayed price response or risk premium? *Journal of Accounting Research* 27, 1-48.
- [12] 3 13. Bernard, Victor L., and Jacob K. Thomas, 1990, Evidence that stock prices do not fully reflect the implications of current earnings for future earnings, *Journal of Accounting and Economics* 13, 305-341.
- [13] 4 14. Boni, Leslie, and Kent L. Womack, 2006, Analysts, industries and price momentum, *Journal of Financial and Quantitative Analysis* 41, 85-109.
- [14] 5 15. Campbell, Sean D. and Steven A. Sharpe, 2008, Anchoring bias in consensus forecasts and its effect on market prices, *Journal of Financial and Quantitative Analysis*, forthcoming 16 16. Chan, Louis K.C., Narasimhan Jegadeesh, and Josef Lakonishok, 1996, Momentum strategies, *Journal of Finance* 51, 1681-1713.
- [15] 7 17. Chen Qu, Liming Zhou, Yue-jia Luo, 2008, "Electrophysiological correlates of adjustment process in anchoring effects", *Neuroscience Letters* 445 (2008) 199-203

Email: Farzaneh.anvari.r@gmail.com

Analysis of lapse risk with censored data: Mellat Insurance case study

Mohammad Azardbad¹

Shahid Chamran University of Ahvaz, Ahvaz, Iran

Khaled Masoumifard

Mellat Insurance, Tehran, Iran

Soroush Amirhashchi

Mellat Insurance, Tehran, Iran

Reyhaneh Jannati

Mellat Insurance, Tehran, Iran

Abstract

Life insurance is a dynamic and attractive market. In fact, everyone who buys life insurance goes through a stochastic process with several stopping times like lapse time and death time. This stochastic process depends on many factors such as gender, profession, age, Time elapsed time of the beginning of the contract, economic parameters and etc. In this paper, we present a approach for generating an internal parametric path process based on the maximum likelihood of exponentiation log-logistic geometric distribution. By using this internal path in the nested simulation approach we can project the cash flow for customers.

Keywords: Lapse, Cash-Flow-Projection, Salvency Capital Requirement

AMS Mathematical Subject Classification [2018]: C52, C53

1 Introduction

Two sets of stochastic scenarios have to be produced: outer scenarios for the evolution of all the variables during first year and inner scenarios for the expectation evaluation. outer scenarios are scenarios where the parameters are obtained from the observed data. In order to be able to simulate the cash-flow for a customer, we need to have detailed information about the parameters of our model. One of the most important risks that we need its information is lapse processes. Two sets of stochastic scenarios have to be produced: outer scenarios for the evolution of all the variables during first year and inner scenarios for the expectation evaluation. Outer scenarios are scenarios where the parameters are obtained from the observed data. In order to be able to simulate the cash-flow for a customer, we need to have detailed information about the parameters of our model. One of the most important risks that we need its information is lapse processes.

¹speaker

2 Main result

In this section, we investigate the lapse probability for policyholders by using a real dataset. In this dataset, every policyholder has either lapsed or is still active. We have used a very general and flexible probability distribution for the lapse time,

$$f(t) = \begin{cases} p_0 & t = 0 \\ (1 - p_0) \cdots (1 - p_{t-1})p_t & 0 < t \leq k \\ (1 - p_0) \cdots (1 - p_{k-1})(1 - p_k)^{t-k}p_k & t > k \end{cases} \quad (1)$$

where $p_0, p_1, \dots, p_k \in [0, 1]$. We can mix this model with a regression model with m independent variables as follows

$$p_j = g(a_j^{(0)} + a_j^{(1)}x_1 + a_j^{(2)}x_2 + \cdots + a_j^{(m)}x_m), \quad j = 0, 1, \dots, k.$$

Where g can be any probability distribution function (e.g. logistic, probit,...). By considering the right-censored data, the likelihood function for this data as follows

$$L(a_0^{(0)}, a_0^{(1)}, \dots, a_0^{(m)}, \dots, a_k^{(0)}, a_k^{(1)}, \dots, a_k^{(m)}) = \prod_{i=1}^n (f(t_i))^{\delta_i} (S(t_i))^{1-\delta_i}. \quad (2)$$

Where S is the survival function of T and δ_i is a binary variable that indicates a censored observation. By maximizing the likelihood function, we can find the estimation of p_j as follow

$$\hat{p}_j(x) = g(\hat{a}_j^{(0)} + \hat{a}_j^{(1)}x_1 + \hat{a}_j^{(2)}x_2 + \cdots + \hat{a}_j^{(m)}x_m).$$

Example: For our dataset, we have estimated the parameters for $k = 5$ and the logistic function $g(x) = \frac{1}{1+e^{-x}}$. Using the maximum likelihood method, the following result is obtained

$$\begin{array}{lll} \hat{a}_0^{(0)} = +0.0054 & \hat{a}_0^{(1)} = +0.0014 & \hat{a}_0^{(2)} = -6.5032 \\ \hat{a}_1^{(0)} = -0.0102 & \hat{a}_1^{(1)} = -0.0441 & \hat{a}_1^{(2)} = -0.4162 \\ \hat{a}_2^{(0)} = -0.0018 & \hat{a}_2^{(1)} = -0.0466 & \hat{a}_2^{(2)} = -0.1547 \\ \hat{a}_3^{(0)} = +0.0026 & \hat{a}_3^{(1)} = -0.0415 & \hat{a}_3^{(2)} = -0.5099 \\ \hat{a}_4^{(0)} = -0.0008 & \hat{a}_4^{(1)} = -0.0036 & \hat{a}_4^{(2)} = -1.4216 \\ \hat{a}_5^{(0)} = -0.0051 & \hat{a}_5^{(1)} = +0.0117 & \hat{a}_5^{(2)} = -0.7452. \end{array}$$

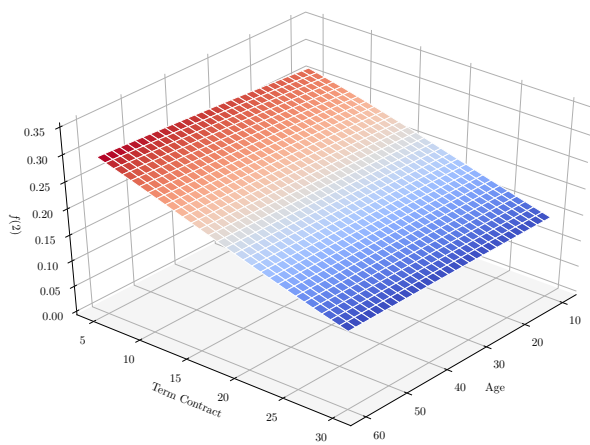
Since we have a lot of censored observations, we expect that the mean of T is bigger than the mean of the real data. For example, if age = 30 and term contract = 15, mathematical expectation of T is 3.84 and mean of data (ignoring the fact that some of them are censored) is about 3.4. In figure 1 we see how the expectation value changes as age and term contract change.

Acknowledgment

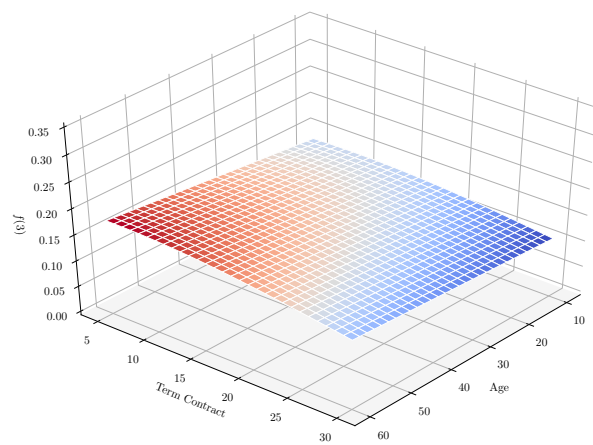
This was supported by the Mellat Insurance Company.

References

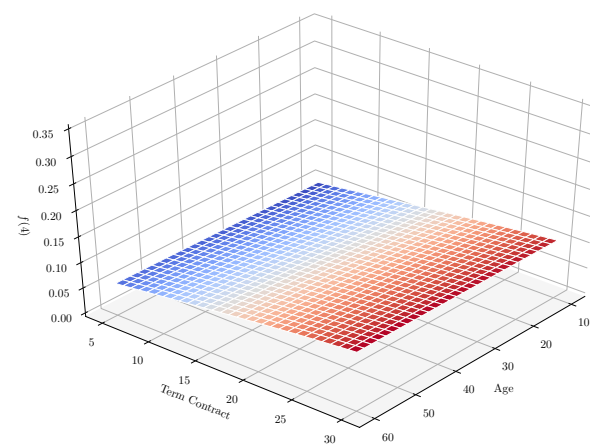
- [1] X. S. Lin and X. Liu, *Life insurance lapse behaviour: evidence from China*, The Geneva Papers on Risk and Insurance-Issues and Practicel 11, no. 4 (2019), pp. 653-678.



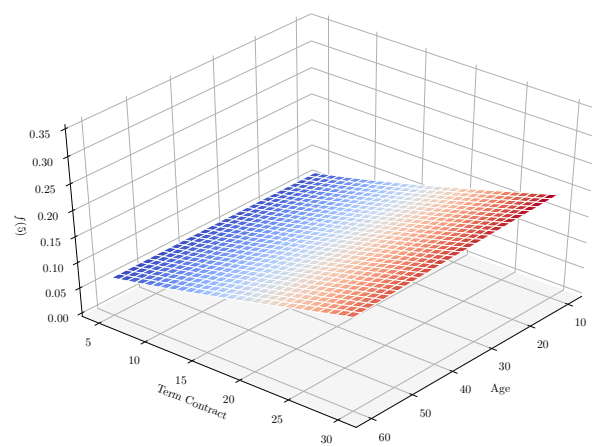
(a) $P(T = 1)$



(b) $P(T = 2)$



(c) $P(T = 3)$



(d) $P(T = 4)$

Figure 1: Probability

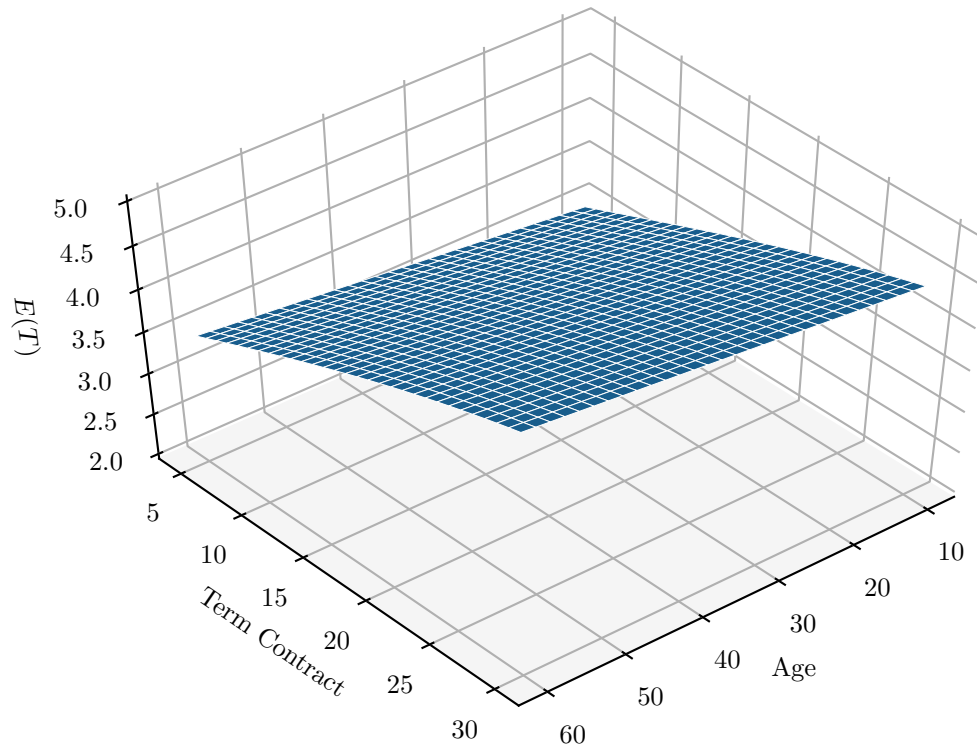


Figure 2: Mean lapse time

e-mail: m.azarbad@hotmail.com
e-mail: k_masoumifard@sbu.ac.ir
e-mail: soroush.amirhashchi@gmail.com
e-mail: R.jannati.k@gmail.com

A simulation-based internal solvency model for determining premium risk of car insurances in Saman Insurance Company

Zahra Barzegar¹
Saman Insurance Company, Tehran, Iran
Shima Ara
Saman Insurance Company, Tehran, Iran

Abstract

Solvency is a concept that determines the ability of insurance companies to meet their long-term fixed expenses and to accomplish long-term expansion and growth. Indeed, the solvency of insurance companies will be determined by measuring the risks that threaten their businesses. It needs to calculate the capital requirement to face expected losses. In this paper, we will provide a methodology based on the compound distribution of portfolio aggregate claim amount to determine capital requirement. To this end, it needs to find the distribution of aggregated loss function, which depends on the frequency and severity of paid losses. The modeling procedure includes two steps. At first step the frequency and severity are modeled separately, then by employing bootstrap algorithms, the distribution of total claim sizes are determined. Illustration of this approach has been provided by applying that on car insurance data of Saman insurance company.

Keywords: Solvency, frequency, severity, aggregate claim, capital requirement

References

- [1] J. Alm. *A simulation model for calculating solvency capital requirements for non-life insurance risk*. Scandinavian Actuarial Journal (2015), (2015): 107-123.
- [2] J. Kočović, J., and M. Koprivica, *An internal model for measuring premium when determining solvency of non-life insurers*. Ekonomski Anali/Economic Annals (2018), 63(217). G. B. Folland, *Real Analysis: Modern Techniques and Their Applications*, 2nd ed., John Wiley, 1999.
- [3] A. Sandstrom, A. *Solvency: Models, Assessment and Regulation*. Boca Raton: Chapman & Hall/CRC (2006).

Email: z.barzegar@samaninsurance.ac.ir

Email: sh.ara@samaninsurance.ac.ir

¹ Speaker

An Early Warning System for Systemic Risk using Breakpoint regression

Saman Ebrahimpour¹
Shahid Beheshti, Tehran, Iran

Zanar Ahmadi
Institute for Advanced Studies in Basic Sciences, Zanjan, Iran

Abstract

Nowadays, systemic risk measures are getting common as a tool for surveillance and regulate the real part of an economy. Although many measures have been introduced in recent years, particularly after the 2007-2008 financial crisis, the main concern is the point that an economic go worse and a crisis may occur to a jurisdiction.

In this article, we tried to introduce a breakpoint regression method to estimate the crisis threshold so that the policy-makers can take precautionary actions before a crisis arises. We considered the value-weighted market Expected Shortfall (ES) as a systemic risk and stock market index as a global indicator. Using breakpoint regression we find some points that potentially could be the threshold of the systemic risk measure.

Keywords: Systemic Risk, Expected Shortfall, Breakpoints, Early Warning.

1 Introduction

Most classical tests against changes in the coefficients of a linear regression model assume that there is just a single change under the alternative or that the timing and the type of change are known. More recently, there has been a surge of interest in recovering the date of a shift if one has occurred or in methods which allow for several shifts at once, see [Bai \(1997\)](#), [Hawkins \(2001\)](#), [Sullivan \(2002\)](#) and [Bai and Perron \(2003\)](#) among many others. In this paper, we are going to introduce some methods for finding breakpoints in regression and time series models.

Consider the following standard regression model:

$$y_i = x_i^T \beta_i + u_i \quad (i = 1, \dots, n)$$

In this section, the regression coefficients are tested for consistency:

$$H_0: \beta_i = \beta_0 \quad (i = 1, \dots, n)$$

In contrast, H_1 assumes that at least one of the coefficients changes over time.

In many cases, it can be assumed that M breakpoint exists where the regression coefficients change. So there is an $M+1$ segment where the regression coefficients are constant and the above model can be rewritten as follows:

$$y_i = x_i^T \beta_j + u_i \quad (i = i_{j-1} + 1, \dots, i_j, \quad j = 1, \dots, m+1)$$

The purpose of this paper is to find a threshold for the expected shortfall (ES) as one of the systematic risk measures.

2 Regression results

In Fig. 1 and Fig. 2, respectively, you can see the time-series graph of the ES values and the total index from 06/01/1390 to 26/11/1395.

In the first method, a simple regression model is considered between ES and total index and the change of regression coefficients is tested.

¹ speaker

To this end, ES is considered as the dependent variable and the Tehran Stock Exchange index as an independent variable, considering the positive correlation between ES and the total index (using Pearson test) using the segments test where the slope The regression line changes can be calculated.

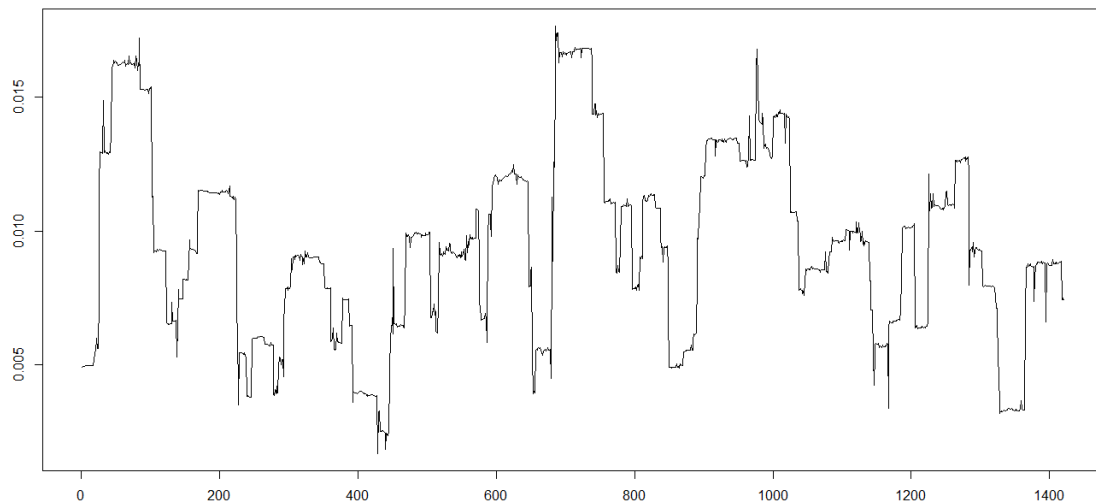


Figure 1: Time series of ES index from 06/01/1390 to 26/11/1395

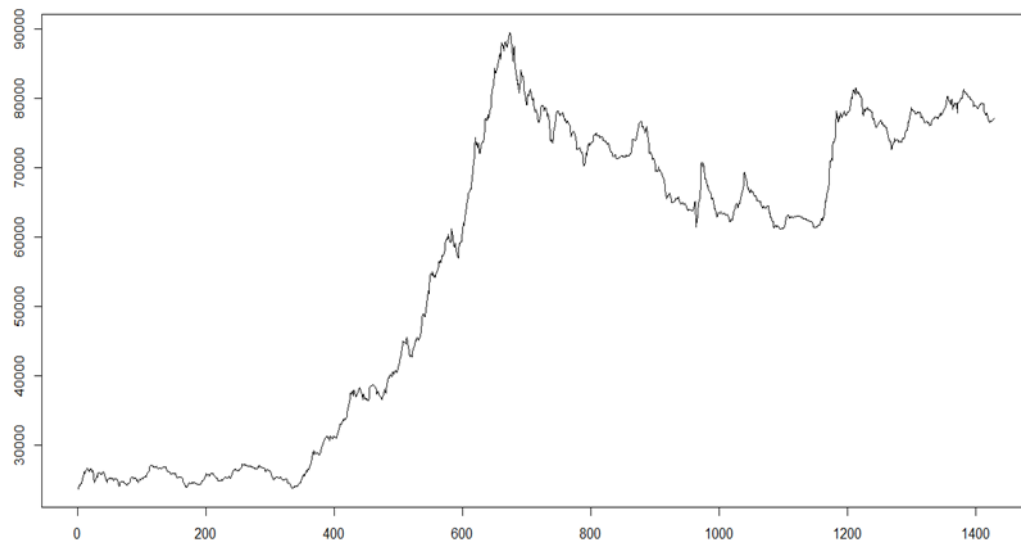


Figure 2: Time series of the total return index from 06/01/1390 to 26/11/1395

In Fig. 3 you can see the point diagram of the total index and the ES in front of each other.

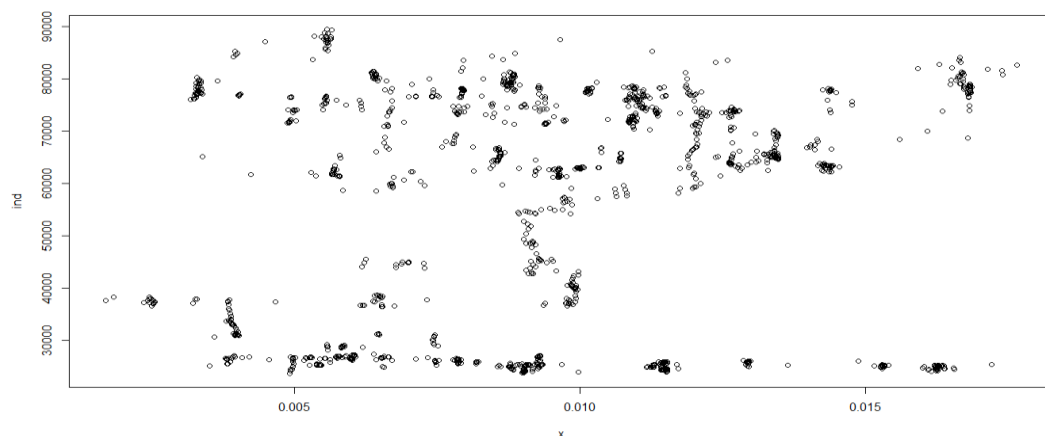


Figure 3: Graph of total index vs. ES

The following table shows the test results.

```

***Regression Model with Segmented Relationship(s)***

Estimated Break-Point(s):
  Est. St.Err
0.016  0.000

Meaningful coefficients of the linear terms:
              Estimate Std. Error t value Pr(>|t|)
(Intercept)    51322      1736   29.572  <2e-16 ***
ES              406331    179735    2.261   0.0239 *
U1.ES           31727202  11516233    2.755      NA
---
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
(Dispersion parameter for gaussian family taken to be 458437121)

Null deviance: 6.7523e+11 on 1428 degrees of freedom
Residual deviance: 6.5327e+11 on 1425 degrees of freedom
AIC: 32560

Convergence attained in 4 iterations with relative change 5.244692e-05

```

As you can see the estimated breakpoint value for ES is .016. This value can be considered as a threshold value of the ES, for the period of 1390 to 1395, the ES value has been calculated for 1421 days, which is 95 days longer than the threshold, meaning that in the 5 years with a probability of .066, the ES value has exceeded the threshold.

2 Time-series results

So far, a breakpoint for ES has been calculated using a simple regression model between index and ES measure. Then, using the time series model, the breakpoints are calculated. For this purpose, we fit a time series model to ES data and calculate the number of breakpoints (points where time-series coefficients change) using Bayesian information criterion tests and least-squares residuals.

The following time series models have been used for this purpose.

$$ES_t = \beta_0 + \sum_{i=1}^{20} \beta_i ES_{t-i}$$

$$ES_t = \beta_0 + \sum_{i=1}^{60} \beta_i ES_{t-i}$$

In Figure 4 you can see the results of both BIC tests.

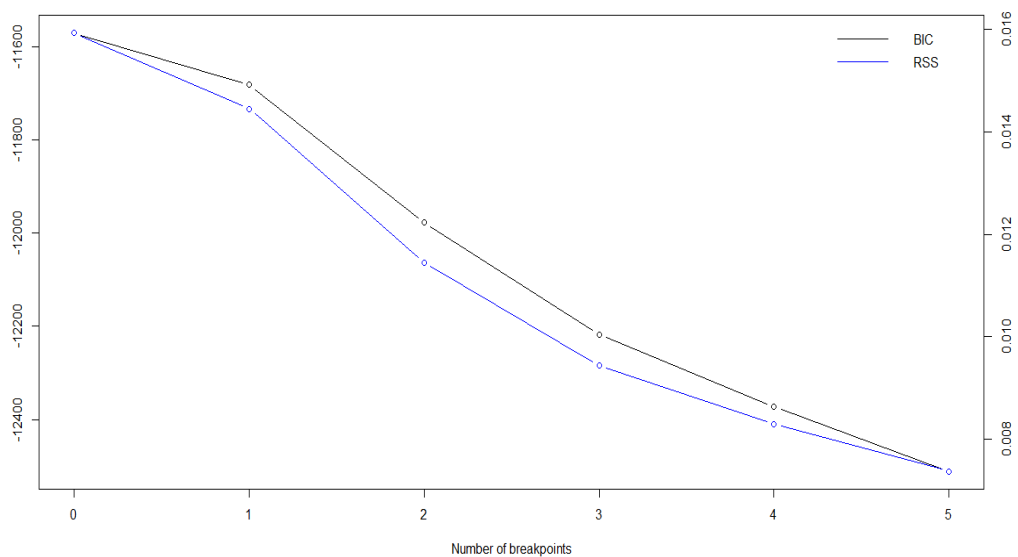


Figure 4: Graph of Bayesian information criterion test and least-squares residuals

As you can see, the five breakpoints give the lowest value in the Bayesian information criterion tests and the least-squares residuals. This means that five test points can be calculated for the test.

Now to calculate the breakpoints, we fit the ES data to AR 20 time series models with lags of 20 and 60. In Figure 5 you can see the graph of ES values over 1429 days (from 01/01/1390 to 26/11/1395) with a lag of 20.

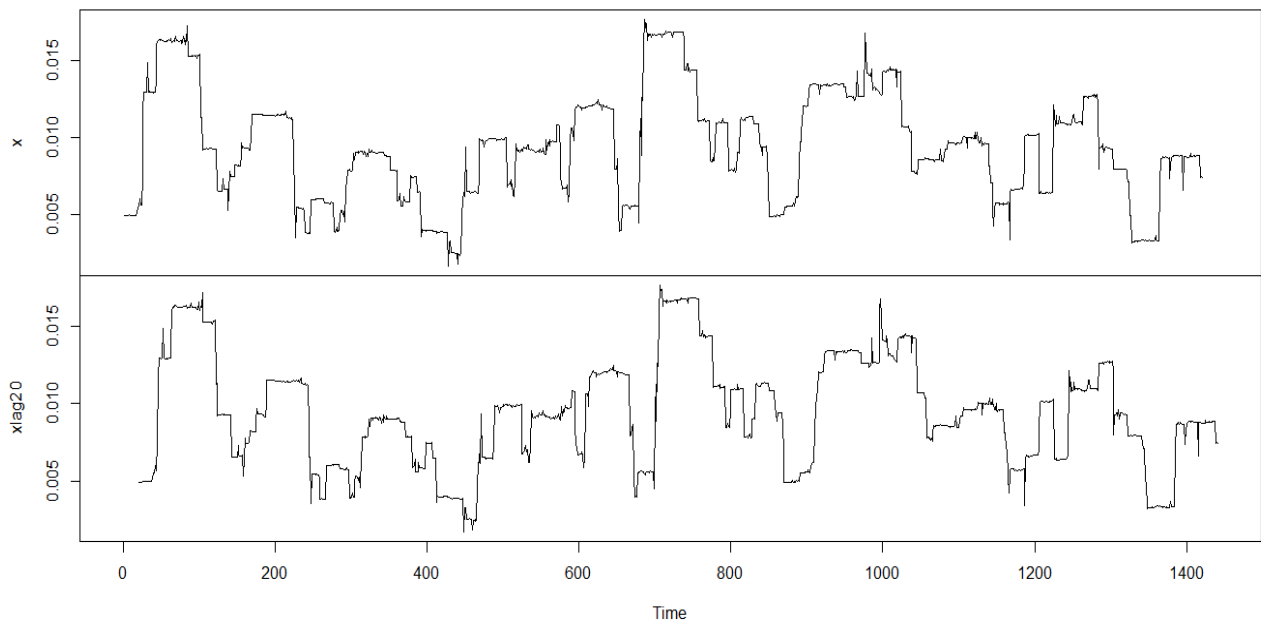


Figure 5: ES time series (top) with 20-day lag ES time series (below)

Using the Breakpoints test, we calculate 5 breakpoints for ES with lag 20.

```
Optimal (m+1)-segment partition:
breakpoints.formula(formula = x ~ xlag22, h = 0.1, breaks = 5,
  data = x)

Breakpoints at observation number:

m = 1    202
m = 2    203 426
m = 3    203 448 680
m = 4    203 448 679    1004
m = 5    203 448 682 869 1009
```

As you can see above, for $m = 5$ breakpoints, points 203, 448, 682, 869 and 1009 are obtained from the data. The points corresponding to these results in ES are .01145, .006231, .0088, .005, and .014393, respectively. In Figure 6 you can see these points on the time-series graph.

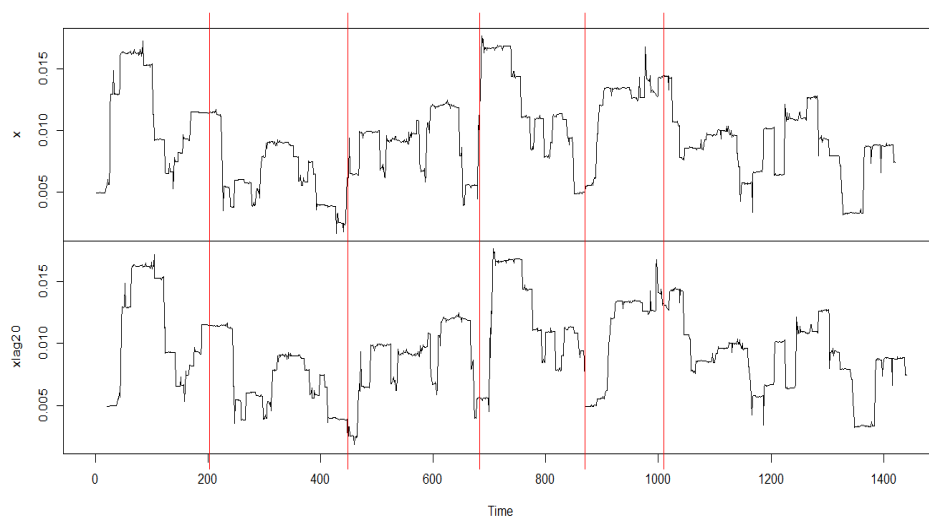


Figure 6: ES time series (top) with 20-day lag ES time series (below) with breakpoints

Since the aim of this report is to find a risk threshold for ES, the results of points .01145 and .014393 can be considered as critical points.

3 Conclusion

Successful implementation of macroprudential policy is contingent on the ability to identify and estimate systemic risk in real-time. Therefore, it is necessary to find a timely risk threshold for timely policymaking. In this paper, methods for finding a threshold of risk measures are introduced and analyzed using Segments Test for Regression and Breakpoints in Time Series. These methods are implemented in the R packages *Strucchange* and *Breakpoints*. In the time series method, we used the BIC method to find the number of breakpoints.

References

- [1] Bai, J., Perron, P., 1998. Estimating and testing linear models with multiple structural changes. *Econometrica* 66, 47–78.
- [2] Arsov, I., Canetti, E., & Kodres, L. (2013). Near Coincident Indicators of Systemic Stress. IMF.
- [3] Zeileis, A., Kleiber, C., & kramer, W. (2003). Testing and dating of structural changes in practice. *Econometrics*. Elsevier.

Email: Saman_ebr@hotmail.com

Email: Zaniara3@gmail.com

Life Table Construction of Active insured members of Social Insurance Fund of Farmers, Villagers and Tribes in 2016

Diba Daraei¹

Shahid Beheshti University, Tehran, Iran

Amin Hassan Zadeh

Shahid Beheshti University, Tehran, Iran

Abstract

Due to the importance of the constructing a life table that provides an accurate estimation of the mortality of the Iranian insured population, this paper is to estimate the mortality rates of insured of the Social Funds of Farmers, Villagers and Tribes in 2016, which actually is a sub-population of 1.5 million people in Iran. In this paper in order to estimate mortality rates of fund, we will only focus on the active insured persons which is roughly 1 million people of total 1.5 million members of fund. In this paper, in order to minimize the effect of data's few-counting of death, the target population has been limited to active insured members in the year 2016. The mortality rates for active insured male members of the fund for ages between 30 and 80 have been estimated. The resulting rates for ages between 30 and 71 have been compared with rates which obtained by insured people of the Social Security Organization, rates of TD (Table de Deces) 88-90 and the rates presented by the United Nations. The results and comparisons show that the estimated rates for the fund are largely reasonable and reliable. It can also be said that the estimated rates for the fund is very close to the rates presented by the United Nations.

Keywords: Mortality rate, the Social Fund of Farmers, Villagers and Tribes, Active Insured members, Makeham mortality model, Social Security Organization, United Nations, Table of TD 88-90

Mathematics Subject Classification [2018]: 62P05

1 Introduction

Demographic rates, such as mortality rates, have an important role in actuarial calculations. Since TD 88-90 Mortality Table, which is prescribed by the Central Insurance of Iran overestimates the mortality rates and is still used in actuarial calculations in Iran, it is important to construct a life table that provides an accurate estimation of the mortality of the Iranian insured population.

In Iran, most of life table construction studies have focused of constructing a life table for the general population. The results of such studies can be found on the Statistical Center of Iran. It can be said the only reliable research conducted for insured people is Eghbal Zadeh and Hassan Zadeh in 2017. [3]

In this paper, we estimate the mortality rates of insured members of the Social Funds of Farmers, Villagers and Tribes (that is called the Tribes Fund or the Fund in this paper) in 2016, which has a population of approximately 1.5 million active insureds in 2018.

The Fund members data in a form of two snapshots of year end 2016 and 2018. The year 2016 data includes 6 items which are Policy number, Date of Birth, Gender, Status, Executive date and Issue date. The year 2018 data, in addition to the items in the year 2016 data, also includes another item named "The number

¹speaker

of premium payment months in 2015” which specifies active insured members in 2015.

The fund has totally 1,509,887 members with 49,833 of whom are in “death” status. Because of high rate of lapse in this fund, it is expected that deaths of some members are not reported, therefore, it is more reasonable to consider only active insured members in the year 2016 as exposure. In order to estimate mortality rates, we consider those who died in the year 2016 and belonged to the active insured members in the year 2016. Following groups have been considered as active insured members in the year 2016:

- Individuals who had premium payment in the year 2015 and did not change their status till the year end 2015 that the number of this group is 920,776 persons.
- Individuals whose issue date were in year 2016 which the number of this group is 92,082 persons.

Considering the two groups above, the fund had 1,012,858 exposure members in 2016. In addition to deaths that have been registered, the fund also had dead members whose deaths had not been registered (therefore their death date are not clear). In order to minimize the effect of data’s few-counting of death, the not-registered deaths should be considered in mortality rates estimation. In the year 2016 the fund has had 6,749 registered-deaths. The number of not-registered deaths of the fund is totally 41,435. This number was obtained from The National Organization of Civil Registrants, but the date of deaths was not available to us. We assumed uniform death distribution for those insureds with issue date after year 2015 and died before year end 2018.

The method which is used for raw mortality rate estimation is as following:

$$\hat{q}_x = \frac{D_x}{E_x + \frac{1}{2}D_x}. \quad (1)$$

In which, $\hat{q}_x$ is mortality rate at age x in the year 2016, D_x is the number of x -aged active insured who were dead in the year 2016 and E_x is the x -aged exposures in the year 2016. See [2] for more details.

For estimating mortality rates, at first D_x and E_x should be calculated. A part of D_x for each age x is determined by intersection between registered-deaths in 2016 and active insured members in 2016. Another part of D_x which is related to not-registered deaths, is determined by intersection between not-registered deaths in 2016 and active insured members in 2016. Since the death date for not-registered deaths are not recorded, not-registered deaths in 2016 according to UDD assumption are determined as follow.

- Intersection between not-registered deaths and the individuals who had premium payment in the year 2015 and did not change their status till the end of year multiplied by $\frac{12}{39}$,
- Intersection between not-registered deaths and Individuals whose issue date were in the year 2016 multiplied by $\frac{6}{27}$.

E_x for each age x is determined by the following:

$$E_x = \frac{1}{12} \times \sum_{i=1}^n (s_{x,i} - t_{x,i}). \quad (2)$$

In which, $s_{x,i}$ and $t_{x,i}$ are the death month and the entrance month number for insured i , respectively. Notice that $s_{x,i} = 12$ for all individuals who did not die till the end of 2016 and $t_{x,i} = 0$ for all individuals who were in the fund as active insured members before 2016. In fact, formula 2 is used to obtain the exact amount of each member exposure. In this paper, it has been considered that all of enters and deaths occurred in the middle of the year 2016.

The exposures for ages between 30 and 80 are such that the estimated mortality rates can be considered reliable. We have decided to choose a minimum exposure 5000. After calculating the crude rates, we have used the Makeham model for smoothing. Based on the Makeham model the force of mortality is given by:

$$\mu_x = A + B \times c^x; \quad 0 < A < 1, \quad 0 < B < 1, \quad c > 1. \quad (3)$$

So the q_x which is the mortality rate for age x is:

$$q_x = 1 - p_x = 1 - e^{-\int_0^1 \mu_{x+t} dt} = 1 - e^{-\int_0^1 (A+B \times c^{x+t}) dt} = 1 - e^{-\left(A + \frac{B(c-1)c^x}{\ln(c)}\right)}. \quad (4)$$

See [1] for more about the Makeham model in mortality.

2 Main results

After calculating the crude rates, we use the least squares method to smooth the rate via the Makeham model. The estimated parameters are as follow.

$$A = 0.0001, \quad B = 0.0001521, \quad c = 1.073.$$

In figure 1 the crude and smoothed rates of Tribe Fund for ages between 30 and 80 are compared.

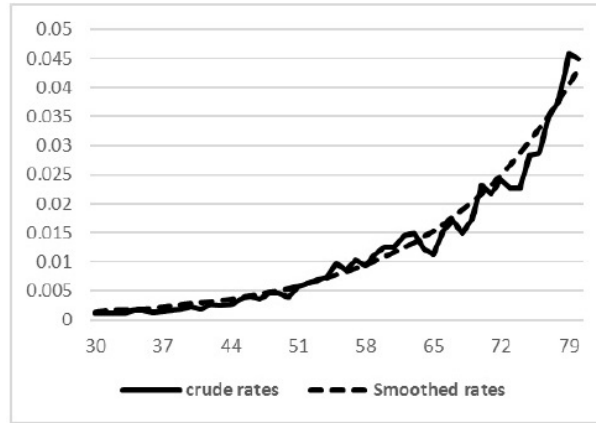


Figure 1: Crude and smoothed rates of Tribes Fund for ages between 30 and 80

In the following, the rates for the fund have been compared with rates from the following three life tables

- Table of TD 88-90,
- The mortality rates of Social Security Organization, [3],
- The rates of United Nations.

The mortality rates from 4 sources mentioned above for ages between 30 and 71 are displayed in the following figure.

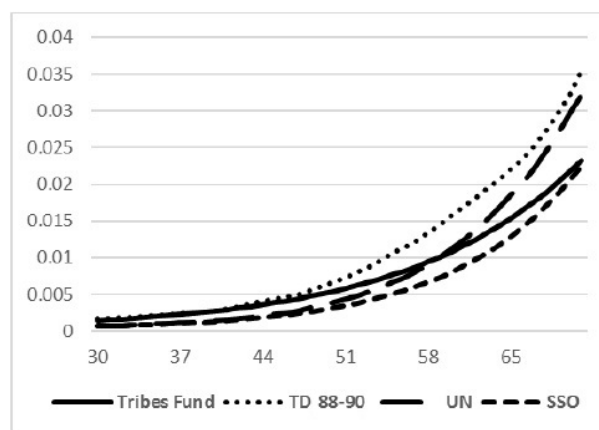


Figure 2: Comparison between rates of 4 sources (Tribes Fund, SSO, UN, TD 88-90) for ages between 30 and 71

The following ratio has been used to compare the rates in pairs.

$$\frac{\sum_{x=m}^n \hat{q}^*}{\sum_{x=m}^n \hat{q}^{**}}. \quad (5)$$

Table 1: The results of comparison between rates of 4 sources for ages between 30 and 71

	Tribes Fund	SSO	UN	TD 88-90
Tribes Fund	1	0.755	1.038	1.375
SSO	1.325	1	1.374	1.822
UN	0.964	0.728	1	1.326
TD 88-90	0.727	0.549	0.754	1

Table 1 illustrates the results obtained by ratio 5.

In Figure 3, the results of another comparison method has been shown, in which, for each age between 30 and 71, the death rate of the mentioned sources are divided by the death rate obtained for the Tribes Fund.

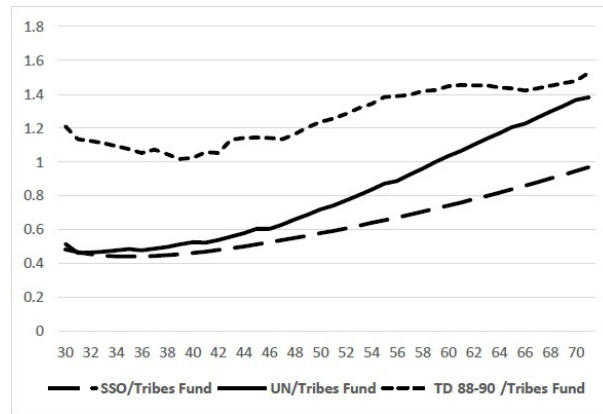


Figure 3: The ratios of the rates of three sources (SSO, UN, TD 88-90) over the Tribes Fund, for ages between 30 and 71

As it can be seen in Table 1, the estimated rates for the Tribes Fund and the United Nations rates are the closest ones. From Figure 3 we have the followings:

- The rates of TD 88-90 are higher than the estimated rates for Tribes Fund for all ages between 30 and 71. However, for ages between 30 and 52 the ratio is close to 1,
- Totally the rates of the Social Security Organization are lower than the estimated rates for Tribes Fund for all ages between 30 and 71,
- For ages between 30 and 60 the estimated rates for Tribes Fund are higher than the United Nations rates and for ages between 60 and 71 this relationship is reversed.

Finally, as the actuaries over the past years have concluded that the mortality rates of TD 88-90 estimate mortality much higher than the actual experience, our research results confirm this fact.

References

- [1] DC. Dickson, M. Hardy, HR, Waters, *Actuarial mathematics for life contingent risks*, Cambridge University Press,(2013).
- [2] HU, Gerber, *Life insurance mathematics*, Springer Science and Business Media, (2013).
- [3] R, Eghbal Zadeh and A, Hassan Zadeh, *Constructing Table of status Switching of insured members of Social Security Organization Using Insured Information (second report)*, Social Security Organization Research Institute, (2017).

Email: D.daraei@sbu.ac.ir

Email: Am_Hassanzadeh@sbu.ac.ir

Modeling and Predicting Health Insurance Claims in Mellat Insurance Company Using A Bayesian Nonparametric Regression Model

Monir Goudarzi¹

Mellat Insurance Company, Tehran, Iran

Soroush Amirhashchi

Mellat Insurance Company, Tehran, Iran

khaled Masoumifard

Mellat Insurance Company, Tehran, Iran

Reyhaneh Jannati Kashani

Mellat Insurance Company, Tehran, Iran

Abstract

Accurate prediction of future claims is a crucial problem in health insurance. In this paper, we fit a Bayesian Nonparametric regression model to health insurance claims data provided by Mellat insurance company. This approach provides model with great flexibility that can accommodate complex characteristics of the regression error distribution such as multimodality, heavy tails, and skewness. The hierarchical structure of the model has the advantage that the parameter estimation is simplified via MCMC methods. The results demonstrate that the fitted model can be improve the predictive accuracy of claims. Also, the model can be used to estimate risk management measures such as VaR and TVaR for the data.

Keywords: Bayesian nonparametric regression, Dirichlet process prior, Predicting, Health insurance claims.

AMS Mathematical Subject Classification [2018]: 62-XX, 62P05

1 Introduction

One of the most important problems in health insurance is modeling and predicting of future claims from policyholders in different risk classes based on past observations of claims made by policyholders in these risk classes.

Nonparametric Bayesian approach is a powerful tool for capturing complex characteristics of the distribution of insurance data such as heavy tails, skewness, or even multimodality. This method place prior distributions on spaces of distributions rather than on parameters of a parametrically specified distribution.

In this paper, we consider a flexible Bayesian nonparametric regression model (Richardson and Hartman, 2018), in particular, a Dirichlet process mixture of log-normals in modeling health insurance claims data provided by Mellat insurance company.

¹speaker

The fitted model will lead to more accurate predictions for claims that can be used for pricing. This model will also be valuable in risk assessment of future obligations.

The rest of this paper is organized as follows: The Dirichlet process and Dirichlet process mixture are briefly described in Section 2. In Section 3, the mathematical structure of the Bayesian Nonparametric regression model is specified. The data analysis is shown in Section 4.

2 Dirichlet Process and Dirichlet Process Mixture

The Dirichlet Process (DP) is most easily characterized by the Sethuraman (1994) construction of it. Let G_0 be a known distribution and let $\alpha > 0$ be a positive constant. Then we say $G \sim DP(G_0, \alpha)$ provided

$$G(\cdot) = \sum_{l=1}^{\infty} w_l \delta_{v_l}(\cdot), v_l \stackrel{iid}{\sim} G_0,$$

where $\delta_v(\cdot)$ defines point mass at v and where $w_l = \xi_l \prod_{i=1}^{l-1} (1 - \xi_i)$ with $\xi_l \stackrel{iid}{\sim} Beta(1, \alpha)$. Therefore, G is a random distribution that is discrete with probability one. G_0 is called the base or centering distribution since $E(G) = G_0$.

The Dirichlet Process Mixture (DPM) takes advantage of the discreteness of the DP. Consider a parametric density function that depends on parameters ν , $f(\cdot | \nu)$, and $\nu | G \sim G$, $G \sim DP(G_0, \alpha)$. We obtain the DP mixture using

$$f(\cdot | G) = \int f(\cdot | \nu) G(d\nu) = \sum_{l=1}^{\infty} w_l f(\cdot | \nu_l),$$

the Sethuraman (1994) construction.

3 Bayesian Nonparametric (BNP) regression model

Let y_i is the log-claim per day of exposure for each policyholder. The Bayesian Nonparametric regression model for is expressed as (Richardson and Hartman, 2018):

$$\begin{aligned} y_i &\sim f(y_i), \\ f(y_i) &= \sum_{l=1}^{\infty} w_l N(y_i | \mathbf{z}' \boldsymbol{\beta}_l, \sigma_l^2), i = 1, \dots, n, \\ (\boldsymbol{\beta}_l, \sigma_l^2) &\sim N(\boldsymbol{\mu}_\beta, \boldsymbol{\Sigma}_\beta) \times IG(a_\sigma, b_\sigma), \\ w_l &= \xi_l \prod_{i=1}^{l-1} (1 - \xi_i), \xi_l \stackrel{iid}{\sim} Beta(1, \alpha), l = 1, 2, \dots \\ \alpha &\sim G(a_\alpha, b_\alpha), \end{aligned}$$

where G , N , and IG represent Gamma, Normal, and inverse gamma distributions, respectively and $\mathbf{z}'$ is a vector of covariates of the model, gender and age.

4 Analysis of the claims data

4.1 The data

The data set is taken from a group health insurance contract in Mellat Insurance Company for 1397. We have claims information on 5461 policyholder that were part of this group plan. The claims were total yearly

costs divided by the number of days of exposure. Each record has two covariates, age and gender that used as factor variables with 2 and 10 levels, respectively.

4.2 Estimation

We analyzed the data by fitting the Bayesian Nonparametric regression model described in Section 3. To estimate the parameters with the Bayesian simulation method, we ran 70000 iterations using the OpenBUGS software. The first 50000 iterations, set as the burn-in period were discarded. To reduce auto-correlation problem, we considered every 20th iteration of chain. The convergence of the MCMC chain was monitored using trace plots, autocorrelation plots, and MC errors of estimates.

The posterior predictive distribution for 6 of covariate combinations are shown in Figure 1.

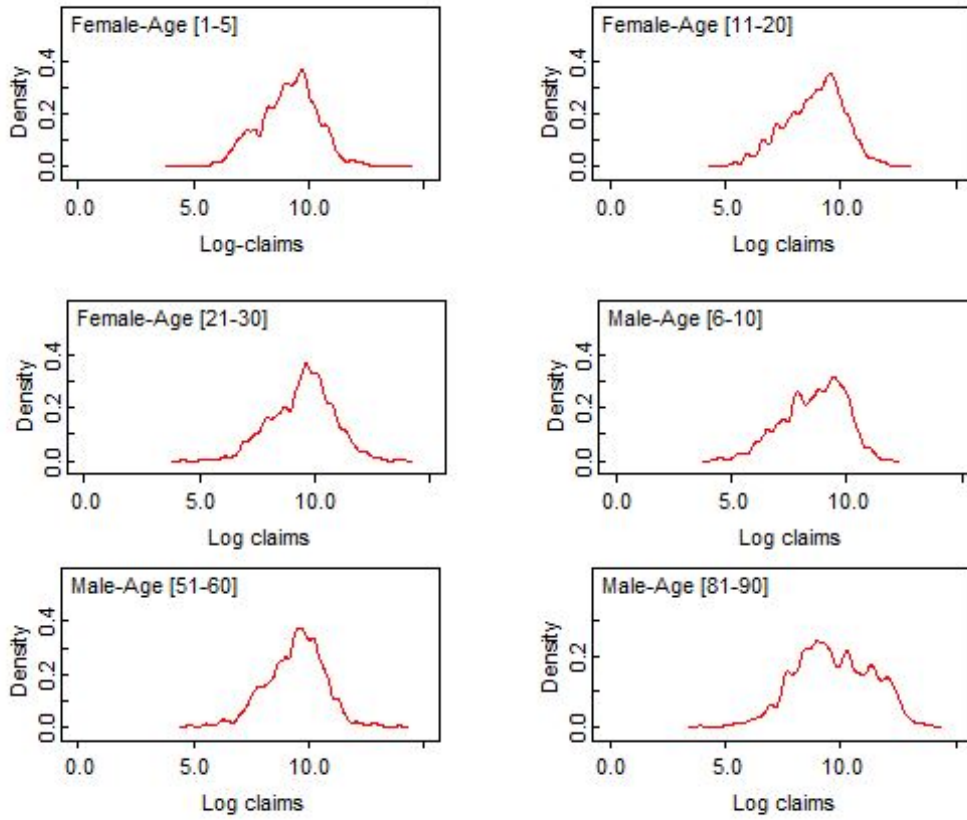


Figure 1: The posterior predictive distribution of the log claims for a number of covariate combinations.

4.3 Risk measures

Following the notation by Klugman et al. (2012), let X denote arandom variable and π_p is the $100p$ quantile of the distribution of X . The Value-at-Risk for a random variable X , denoted as $VaR_p(X)$, is the same as π_p and $P(X > \pi_p) = 1 - p$.

Also, for a random variable X , the Tail-Value-at-Risk, denoted as $TVaR_p(X)$, is the conditional expectation of X given that X exceeded the $100p$ quantile of the distribution, i.e.,

$$TVaR_p(X) = E(X | X > \pi_p).$$

Table 1 summarizes the results of the risk measures for a number of covariate combinations.

Table 1: Summary of risk measures.

	Median	Posterior credible intervals	$VaR(0.95)$	$TVaR(0.95)$
Male-Age [6-10]	8.75	[5.842-10.8]	10.5	10.882
Female-Age [21-30]	9.61	[6.884-11.92]	11.5	12.068
Male-Age [81-90]	9.574	[6.568-12.62]	12.32	12.747

References

- [1] G. W., Fellingham, A. Kottas, and B. M. Hartman, *Bayesian nonparametric predictive modeling of group health claims*, Insurance: Mathematics and Economics, 60 (2015), pp. 1-10.
- [2] S. A., Klugman, H. H., Panjer, and G. E. Willmot, *Loss Models: From Data to Decisions*, fourth ed. John Wiley Sons, Hoboken, NJ, 2012.
- [3] R., Richardson, and B., Hartman, *Bayesian nonparametric regression models for modeling and predicting*, Insurance: Mathematics and Economics, 2018.
- [4] J., Sethuraman, *A constructive definition of dirichlet priors*, Statistica sinica, 1994, pp. 639-650.

e-mail: Goudarzi.monir@gmail.com

e-mail: soroush.amirhashchi@gmail.com

e-mail: K.masoumifard@sbu.ac.ir

e-mail: r.janati@mellatinsurance.com

Dynamic Measurement of Iran Interbank Network Stability

Tayebeh Zanganeh

Islamic Azad University, Science and Research Branch, Tehran, Iran

Mohammad Ali Rastegar

Tarbiat Modares University, Tehran, Iran

Abstract

Banks' commitments to each other mainly arise in the interbank market, which can lead to increased systemic risk due to the spillover effect. Therefore, the objective of this paper is to analyze the network dynamic stability of the Iranian overnight money market through methods of statistical mechanics applied to complex networks. The results show that the network structure changes during time depending economic conditions. Systemic risk measures such as clustering coefficient, average short path, heterogeneity and centrality, show that the networks systemic risk increases and then by occurring default and crisis in one bank, default spillover during the domino effect in whole network. Also, in the event of failure, the most vulnerable group is to privatized and specialist governmental banks, and the private banks, due to the high volume of exchanges and net negative flows, can put a considerable systemic risk to the interbank market network. Moreover, the signals of speculative activity by private banks are found.

Keywords: Network stability, systemic risk, clustering coefficient, heterogeneity, centrality

Mathematics Subject Classification [2018]: 54A06

1 Introduction

After the innovative researches on small-world networks by Watts and Strogatz and scale-free networks by Barabasi and Albert, the study of complex networks has received increasing attention. Complex networks have become a general method for studying common properties of complex systems in the real world, and have penetrated into statistical physics, social sciences, biological sciences and many other fields. Applications of network theory in economic systems can be useful in considering explicitly the relations among economic agents. Many empirical analyses of economic systems have been constructed with the application of network tools, such as the world trade web, e-commerce, the correlation network of stock price returns and commercial credit among firms, financial credit from banks to firms and interbank credit.

In the banking system, an intricate web of claims and obligations links the balance sheets of a wide variety of intermediaries, such as banks and hedge funds, into the structure of a network. As for the banking system, there is abundant theoretical economic literature on contagion risk and systemic risk which suggest various topological structures of the banking system, such as the complete and incomplete interbank structures, the 2-D directed lattice, money-center structure, a random network, and so on.

In addition, some literature drew the conclusion that the banking system stability depended on its structure.

Therefore, it is of great importance to understand the structure of the real banking system. However, the structure of the real banking system is not completely in accord with the theoretical structure those scholars postulated.

Interbank market permits liquidity exchanges among financial institutions through facilitating the allocation of the liquidity surplus to illiquid banks. Complex network relationships are formed through interbank lending, payment and settlement, discount and guarantee. On one hand, the complex debtor-creditor relationships between banks provide channels for interbank liquidity exchanges, but on the other hand, they also become potential paths for financial contagion, which may trigger the domino effect. For example, the US sub-prime mortgage crisis broke out in August 2007, resulting in large number of banks failed (such as Lehman Brothers, Washington Mutual Bank, Colombia Trust, etc), which quickly evolved to a global financial crisis and greatly damaged the global financial system.

It is well known that the topology of a network (for example the Internet connectivity map, the World Wide Web, author collaboration networks, biological networks, communication networks, power networks) affects its functionality and stability.

Our paper focuses on the network analysis of the overnight maturity on the market for interbank deposits of Iran. Our data set is composed of monthly banks operating in the Italian market between 2010-2015 consisting 66 adjacency lending matrixes. For every month of trading we compute the network of debts and loans. The main objectives are to understand the network topology of Iranian interbank market.

2 Main result

The networks parameters is shown in the below table. By analyzing some of these measures we can determine the systemic risk and instability of interbank network.

kurtosis	skewness	STD	Max	Min	medium	mean	parameters
0.683	-0.670	4.921	28.000	4.000	19.500	19.106	number of nodes
-0.543	0.612	49.436	205.000	3.000	63.000	74.697	number of links
-1.042	0.110	0.098	0.533	0.167	0.342	0.337	network density(undirected)(connectivity)
-0.929	0.259	1.732	7.321	0.667	3.472	3.523	average degree of a node
-0.337	-0.305	0.340	2.430	1.000	1.760	1.737	average path length
-0.427	-0.078	1.445	7.000	1.000	4.000	4.182	network diameters
1.139	1.387	0.527	3.000	1.000	1.000	1.318	network radius
-0.531	-0.318	0.118	0.456	0.000	0.246	0.232	clustering coefficient
-0.535	0.062	0.095	0.526	0.078	0.311	0.338	Network centralization(undirected)
-0.848	0.419	3.167	13.143	1.333	6.000	6.424	average numebr of neighbors
-1.129	0.013	0.094	0.680	0.354	0.519	0.512	network heterogeneity(undirected)
-0.556	0.142	0.892	4.412	0.500	2.105	2.105	Average of Eccentricity

Systemic risk measures such as clustering coefficient, average short path, heterogeneity and centrality, show that the networks systemic risk increases and then by occurring default and crisis in one bank, default spillover during the domino effect in whole network. Also, in the event of failure, the most vulnerable group is to privatized and specialist governmental banks, and the private banks, due to the high volume of exchanges and net negative flows, can put a considerable systemic risk to the interbank market network. Moreover, the signals of speculative activity by private banks are found.

Average is around 2 which shows that the interbank network is small world network. As much the average short path and clustering coefficient is small the network gets more stable. This index has increased steadily since 1392, therefore the network's stability reduces.

The higher the clustering coefficient, the easier it is for a bank to transmit to a neighbor, and thus the more systemic risk is to the other banks.

The structure of Iran's interbank network has become more vulnerable to crisis and default from these measures point s of view.

Another important concept of the network is the network centrality, which is the extent to which a network has one or more specific nodes that other nodes are clustered around.

increasing the heterogeneity and centrality of the network makes the network more similar to the star networks which has more systemic risk.

References

- [1] Aldasoro, I., Gatti, D. D., & Faia, E. (2017). Bank networks: Contagion, systemic risk and prudential policy. *Journal of Economic Behavior & Organization*, 142, 164-188.
- [2] Di Gangi, D., Sardo, D., Macchiati, V., Minh, T. P., Pinotti, F., Ramadiah, A., & Cimini, G. (2018). Network Sensitivity of Systemic Risk. *arXiv preprint arXiv:1805.04325*.
- [3] Erol, S., & Vohra, R. (2018). Network formation and systemic risk. Available at SSRN 2546310.
- [4] Engel, J., Pagano, A., & Scherer, M. (2019). Reconstructing the topology of financial networks from degree distributions and reciprocity. *Journal of Multivariate Analysis*.
- [5] Krause, S. M., Štefančić, H., Zlatić, V., & Caldarelli, G. (2019). Controlling systemic risk-network structures that minimize it and node properties to calculate it. *arXiv preprint arXiv:1902.08483*.
- [6] Leventides, J., Loukaki, K., & Papavassiliou, V. G. (2019). Simulating financial contagion dynamics in random interbank networks. *Journal of Economic Behavior & Organization*, 158, 500-525.
- [7] Armelius, Hanna, Christoph Bertsch, Isaiah Hull, and Xin Zhang. "Spread the Word: International Spillovers from Central Bank Communication." *Journal of International Money and Finance* (2019): 102116.
- [8] Tonzer, Lena. "Cross-border interbank networks, banking risk and contagion." *Journal of Financial Stability* 18 (2015): 19-32.

Email: t_zangene@yahoo.com

Email: m_rastegar@yahoo.com

Valuation of the new health product with a popular rider and limited benefits

Sara sadat Moosavi¹

Shahid Beheshti University, Tehran, Iran

Amir T. Payandeh Najafabadi

Shahid Beheshti University, Tehran, Iran

Abstract

This paper introduces a new health product for first time which provides the benefits of insurer and policyholder. This innovative product includes the benefits of long-term care, limited hospitalization coverage and guaranteed lifelong withdrawal benefit option. Besides, the premium of this product is directly affected by the return of financial market and this product protect the benefits of policyholder against the downside risk.

Keywords: Variable annuity, Long-term care, Life care annuity, Guaranteed lifetime withdrawal benefit.
Mathematics Subject Classification [2018]: G22 , I1.

1 Introduction

Life Care Annuity (LCA) is a combination of lifetime annuity and long-term care insurance that discussed by Murtaugh et al. [2] for first time. They studied the effect of the positive correlation between mortality and disability on the LCA and proved that this product can reduce the cost of regular LTC.

Insurance companies offer variable annuity products along with a variety of riders, which are called guarantee. These guarantees increase the willingness of policyholder to buy insurance policies. Life Care Annuity-Guaranteed Lifetime Withdrawal Benefit is an example of such new product that suggested by Hsieh et al. ([1]) in 2018 for the first time. In this popular product, LCA and withdrawal benefits are provided together. This article introduce a new LCA-GLWB product that considers limited hospitalization coverage as its benefit. Limited hospitalization coverage has been introduced by Yang et al. ([4]) in 2016. They proposed an evaluation model that can accurately provide fair premiums for limited coverage policies.

This paper is organized as follows. In section 2, the details of the product have been reviewed. Moreover the methods that are employed for evaluating the product are discussed. Numerical results are provided in section 3 and we conclude with a discussion of our findings in Section 4.

2 Product specifications

At the beginning of the contract(State 1 in figure 1), the policyholder pays a lump sum (w_0) to the insurance company. Lamp sum has been invested in an investment fund, grows based on fund's return rate $R(t)$ and reach to $W(t)$ at time t . In this paper, we assume that insurer invest w_0 under a Geometric Brownian Motion (GBM) process. Therefore the fund's return rate $R(t)$ satisfies

$$R_t = \exp\left\{\left(\mu - \frac{\sigma^2}{2}\right) + \sigma(B_t - B_{t-1})\right\}; t = 1, 2, \dots, K_x, \quad (1)$$

¹speaker

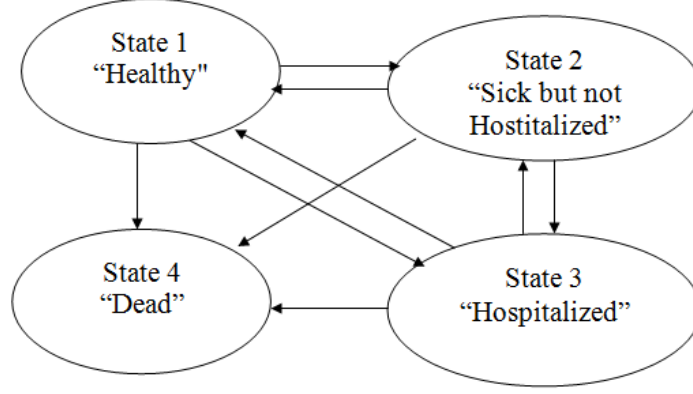


Figure 1: The health status of the policyholder.

where K_x denotes the number of completed future years lived by the policyholder, B_t is the standard Brownian motion process and μ and σ^2 are expected drift rate and volatility of the process, respectively. Moreover at the beginning of year t , the insurer withdraws the fixed management fees M and guaranteed fees αW_t . Limited hospitalization costs (HB_t^x), healthcare costs (SB_t^x) and withdrawal benefits ($g_t = \beta_t w_0$) are other benefits that paid to policyholder at the ending of year t . Suppose that the healthcare benefits, that pay back to policyholder when s/he moves to state 2 at time t , modeled by

$$SB_t^x = cw_0(1 + \pi)^t, \quad (2)$$

where π and c are a fixed inflation rate and positive constant respectively. Now suppose that this product will paid b_t in t 'th year of the policy and will reach to zero after L days. Therefore, if T_t^x denote the total number of days for hospitalization in the t 'th policy year for an x -aged insured, then

$$HB_t^x = b_t T_t^x I(0, D_x) + \left(\sum_{i=1}^{D_x} b_i L - \sum_{i=1}^{t-1} b_i T_i^x \right) I\{D_x\} + 0I(D_x, \infty), \quad (3)$$

where D_x is the first time that total number of hospitalization days is greater than L and $I()$ is an indicator function. Finally the contract will be expired either policyholder's invested amount reaches to zero or s/he touches state 4(death). Based on the given description the account value at year t before and after withdrawals, which is represented by W_t^- and W_t^+ respectively, satisfies

$$\begin{aligned} W_0^- &= w_0, \\ W_0^+ &= \max(W_0^- - \alpha W_0^- - M, 0), \\ W_t^- &= R_t W_{t-1}^+; t = 1, 2, \dots, K_x, \\ W_t^+ &= \max(W_t^- - \alpha W_t^- - M - g_t - HB_t^x - SB_t^x, 0); t = 1, 2, \dots, K_x - 1, \\ W_{K_x}^+ &= \max(g_{K_x} + HB_{K_x}^x + SB_{K_x}^x, W_{K_x}^-). \end{aligned} \quad (4)$$

Now the fair value under an equivalent martingale measure Q is expressed as

$$E_Q\left(\sum_{t=1}^{K_x} \frac{g_t + SB_t^x + HB_t^x}{\beta(t)}\right) + E_Q\left(\frac{\max(0, W_{K_x}^- - (g_{K_x} + SB_{K_x}^x + HB_{K_x}^x))}{\beta(K_x)}\right), \quad (5)$$

where $\beta(t)$ is a risk-less asset that is known as the money market account and $\beta(0) = 1$. In next section, we evaluate proposed product and find the fair value of it.

Table 1: The actuarial values of the product using Monte Carlo method.

Entry age	60	65	70	75
Fair value	332407.7	338223.6	346139.9	348427.8

3 Numerical Results

This section presents the numerical results to estimating the fair value of the proposed product. Assume that fixed daily payment for hospitalization (b_t) and total limited hospitalization days (L) are 1000 and 10 respectively. Moreover, assume that fixed management fees (M), guaranteed fees (α), c , inflation rate (π) and withdrawal benefits (β_t) are 300, 0.008, 6%, 5% and 2% respectively. As defined in the previous section, T_t^x follows a Poisson-Poisson distribution with assuming 0.3 and 12.3 for primary and secondary time-varying intensity respectively. Mean and standard deviation of the invested mutual fund at time t under GBM process are 0.04 and 0.16 respectively and

$$\mu_{x+t}^{ij} = A_{ij} + B_{ij} \exp\{C_{ij}(x - 68.5 + t)\} + D_{ij}(x + t); \forall 65 \leq x + t \leq 120, \quad (6)$$

where

$$A = \begin{bmatrix} 0 & -0.0322 & 0.0096 & -0.0234 \\ 1.0400 & 0 & -0.3380 & 0.0294 \\ 0.1740 & 0.5450 & 0 & 0.1850 \\ 0 & 0 & 0 & 0 \end{bmatrix}, B = \begin{bmatrix} 0 & 0.0519 & 0.0021 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0.0056 \\ 0 & 0 & 0 & 0 \end{bmatrix}, \quad (7)$$

$$C = \begin{bmatrix} 0 & 0.0435 & 0.1741 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0.1330 \\ 0 & 0 & 0 & 0 \end{bmatrix}, D = \begin{bmatrix} 0 & 0 & 0 & 0.004 \\ -0.0113 & 0 & 0.0083 & -0.0002 \\ -0.0015 & -0.0047 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}.$$

Table 1 is presented the actuarial value of the product. It can be seen that the fair value of this product has been increased by age, because older insured are more likely to use healthcare and hospitalize benefits than younger people.

Hospitalization days, withdrawal benefits, inflation rate and coefficient of healthcare costs are important factors to evaluating the product. Since increasing the value of them directly increase the benefits (or annuity) that paid to policyholder, therefore we expected as those increase the price of the product. Figure 2 shows the impacts of these factors on fair value. As we expected, figure 2 shows that the fair value of the product impacted by the value of the hospitalization days, withdrawal benefits, inflation rate and coefficient of healthcare costs, directly.

4 Conclusion

In this paper, we analytically introduce a new health product and evaluate it. An x age policyholder pays w_0 at the beginning of the contract and receives limited hospitalization costs, healthcare costs and withdrawal benefits. Finally, this product expires either the policyholder dies or his/her account value reaches to zero. We believe that this product includes the benefits of guaranteed income streams and LCA for the policyholder and is more attractive for policyholders.

References

- [1] Hsieh, M.h., Wang, J. L. , Chiu, Y.F., and Chen, Y.C., *Valuation of variable long-term care annuities with guaranteed lifetime withdrawal benefits: A variance reduction approach*, Insurance: Mathematics and Economics . 78 (2018), pp. 246–254.
- [2] Murtaugh, C. M., Spillman, B. C. and Warshawsky, M. J., *In sickness and in health: An annuity approach to financing long-term care and retirement income*, Journal of Risk and Insurance. (2001), pp. 225–253.

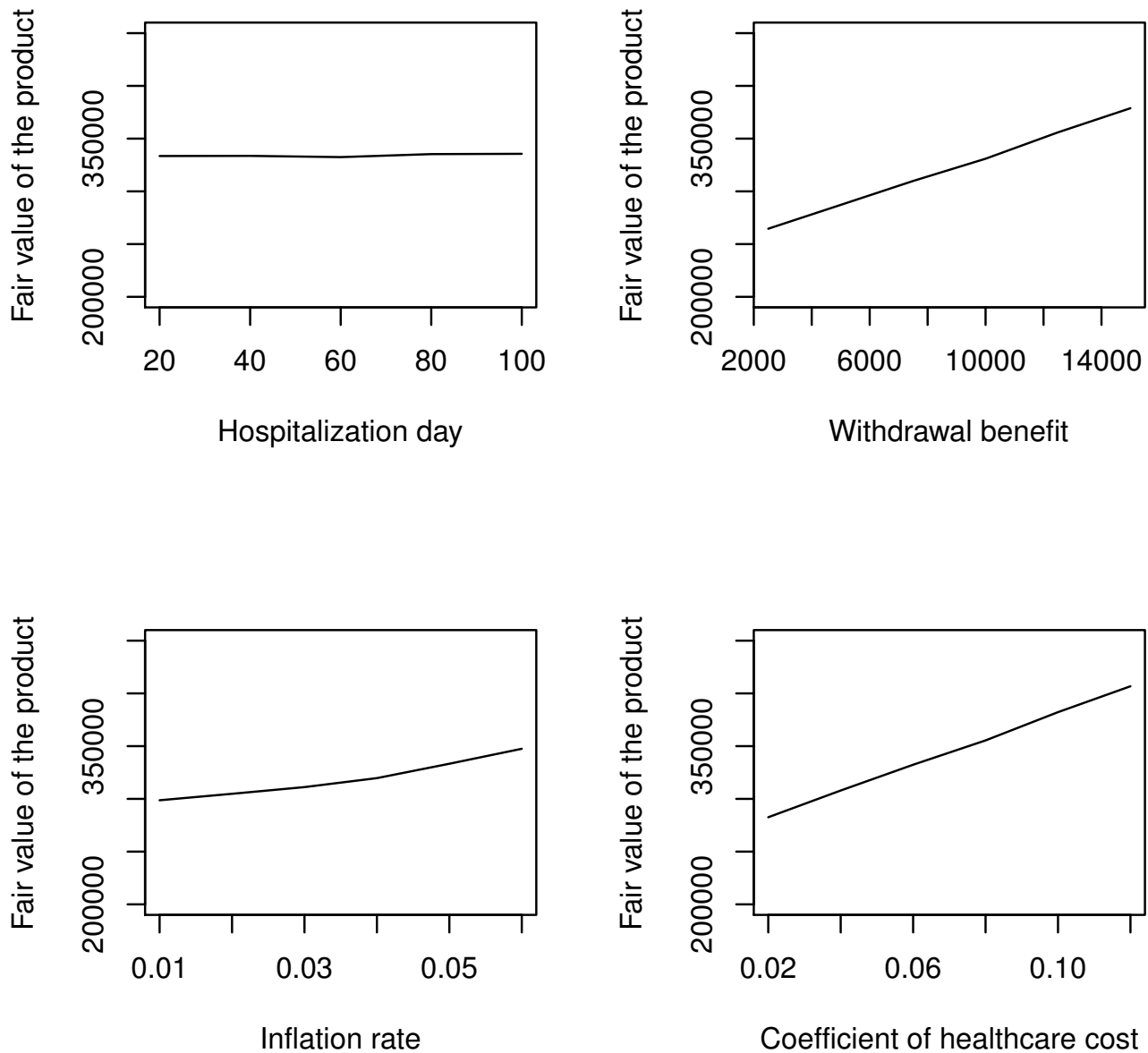


Figure 2: The fair values of the product using Monte Carlo method for 60-aged insured.

- [3] Pitacco, E., *Actuarial models for pricing disability benefits: Towards a unifying approach*, Insurance: Mathematics and Economics. 16(1) (1995), pp. 39–62.
- [4] Yang, S. Y., Wang, C. W. and Huang, H. C., *The valuation of lifetime health insurance policies with limited coverage*, The Journal of risk and insurance. 83(3) (2015), pp. 777–800.

Email: s_ moosavi@sbu.ac.ir

Email: amirtpayandeh@sbu.ac.ir

Mean Field Games in Finance

Erfan Salavati¹

Amirkabir University of Technology, Tehran, Iran

Abstract

The modern mathematical finance starts with the idea of arbitrage pricing. The origin of this idea goes back to the works of Arrow and Debreu on general equilibrium of markets with a representative agent.

Since then, mathematical finance has been mainly concerned with the derivatives pricing, portfolio optimization and risk quantification. These problems, although interesting from a mathematical point of view, has lost their connection to micro-economic foundation of mathematical finance, and there has not been much studies on the equilibrium of financial markets.

In recent years, the new emerging topic of Mean Field Games (MFG) has created proper framework for the study of equilibrium in financial markets.

MFG is a model for strategic decision making of N agents when N is a large number. Each agent has a set of actions A . All agents are utility maximizing with a utility function $J(\alpha^i, \bar{\alpha})$ where α^i is the i th agent's action and $\bar{\alpha}$ is an overall average of the actions of the other agents. Hence the average $\bar{\alpha}$ is the factor that makes the interaction of the agents. The main question of interest in MFG is finding Nash equilibrium of this strategic situation when $N \rightarrow \infty$.

MFG has been successfully applied to several problems in finance such as *Modelling the Bank Runs*, *Systemic Risk*, *Price Impact* and *Limit Order Books*.

In this talk we study a limit order market consisting of inhomogeneous agents and show that how the Nash equilibrium gives rise to a Forward-Backward Stochastic Differential Equation (FBSDE) and discuss some of the consequences.

Keywords: Mean Field Game, Differential Game, Limit Order Market

AMS Mathematical Subject Classification [2018]: 91A13, 91A23, 91G80, 93E20

References

- [1] Carmona, Ren, and Francois Delarue. Probabilistic Theory of Mean Field Games with Applications I-II. Springer Nature, 2018.
- [2] Casgrain, Philippe, and Sebastian Jaimungal. "Mean-field games with differing beliefs for algorithmic trading." arXiv preprint arXiv:1810.06101 (2018).
- [3] Cardaliaguet, Pierre, et al. The Master Equation and the Convergence Problem in Mean Field Games:(AMS-201). Vol. 381. Princeton University Press, 2019.

e-mail: erfan.salavati@aut.ac.ir

¹speaker

Factor investing meets predictability

Ehsan Adman¹

Isfahan University, Isfahan, Iran

Mahmoud Botshekan²

Isfahan University, Isfahan, Iran

Abstract

This paper examines if the predictability of factors' premia can improve the performance of portfolio optimization techniques aim in optimal investing in these factors. We apply different methods of optimization to the well-known factors used in asset pricing models, including market factor, SMB, HML, RMW, CMA and momentum while trying to predict their premiums using different approaches of predictability in order to achieve superior investment performance, compared to when using their historical means as their predictions. We used different linear and non-linear methods to predict factor premiums in various time frames. We found that using predictability of factor premiums can improve the investment performance in factors and this improvement increase as we widen the time frame of the predictions.

Keywords: Factor Investing, Predictability, Factor Premiums, Non-linear Models

Mathematics Subject Classification [2018]:

1 Introduction

Factor investing is a new approach emerged in recent years which introduce new asset classes based on risk factors developed in multifactor asset pricing models as new assets for investing. The main idea stems from asset pricing literature which introduce new systematic risk factors like size, value and momentum in the form of hedging portfolios which capture part of risk premia not captured by traditional market risk premium. The real investment in these factors can be done through investing in a representative index or mutual funds formed based on the above mentioned risk attributes.

Like other assets, portfolio optimization techniques like the mean-variance framework of Markowitz (1952) can be used to determine optimal weights for investing in these factors. But as Demiguel et al (2009) among many others have shown the main obstacle for implementing these techniques is finding a good estimate for expected return of the assets. One the other hand, there are a vast literature on predictability of market factor, but the predictability of other factors' premia has been rarely studied in the literature.

Our first motivation in this paper is to examine predictability of factor risk premia like size, value, momentum, investment, and profitability which are the most prominent factors in asset pricing models and can be used in factor investing context. The second motivation is to test if these evidence of predictability can improve the out-of-sample performance of portfolio optimization techniques compare to methods that just use historical average as the expected return. In coming sections, we explain various methods based on predictive regressions we used to examine and improve the predictability of factors' premia, the various predictive variables we used, and also various optimization techniques we employed to test economic significance of the existing predictability evidence.

2 Method

2.1. Predictability of factor premiums

¹ PHD student in financial engineering, University of Isfahan

² Assistant Professor, University of Isfahan

To examine predictability, data starts from July 1963 to March 2018 which consists of 657 months. We use a rolling window approach to estimate and then to forecast one-step-ahead premia. The rolling window is 240 months for all methods. The factors used are Market factor, SMB, HML, RMW, CMA and momentum. The first 5 factors are Fama-French (2015) factors and Carhart (1997) momentum factor.

To investigate the improvement of investment performance by using factors' predictability we apply predictions on time frames of monthly, annually and 2-yearly. In this way the return calculation is always on monthly data but in the case of annual and 2-yearly, first we use annual and 2-yearly predictions respectively to calculate weights and then we calculate the return of the portfolio on monthly returns.

In this paper factor premiums are predicted using a range of linear and non-linear models. Methods are as follows:

2.1.1. Multivariate simple regression:

$$E(f_{j,t+1}) = \hat{a}_f + \sum_{k=1}^3 \hat{a}_k Z_{k,t} + \sum_{l=1}^5 \hat{a}_l C_{l,t} \quad (1)$$

In this regression we use K=3 macro-economic state variables including risk free rate, GDP growth and output gap and L=5 stock market state variables including term spread, default spread, dividend to price ratio, volatility and trend to predict factor premiums.

2.1.2. Multivariate threshold regression:

In this method the same variables as the previous method are used, but here we use term spread as a threshold to divide the sample into two sub-sample and to estimate two different regressions in different states of economy determined by the term spread as a proxy for economy states.

2.1.3. Combination of univariate regressions

This approach employs eight separate univariate regressions to predict each factor premium and then combine these forecasts with an equally weighted average of all forecasts.

2.1.4. Combination of threshold univariate regressions

This method is the same as the previous method in all aspects except that here we use the univariate threshold regressions in which the threshold variable is the same variable used in the univariate regression. So again eight forecast are estimated for each factors and then they are combined with an equally weighted average.

2.1.5. Combination of threshold univariate regressions with positive forecast constraint

In this method the positive forecast constraint of Campbell and Thompson (2007) is use in threshold regressions of previous method in order to achieve more accurate and reasonable forecasts.

2.2. Optimization techniques

This paper investigate whether using predictability improves the performance of various factor investing techniques. we examine the performance of a couple of common optimization models. Models used in this paper are as follows:

2.2.1. Mean-variance optimization, an out-of-sample approach ($\hat{M}\hat{V}$)

This approach is the well-known mean variance optimal portfolio (Markowitz, 1952) in which the weights of optimal portfolio is calculated as follows:

$$\frac{\sum_t^{-1} \mu_t}{1' \sum_t^{-1} \mu_t} \quad (2)$$

In which $\sum_t$ is the covariance matrix of assets and μ_t is the vector of expected return of assets. In this approach both mean and covariance matrix are estimated from the sample data so they may be estimated with error. Once the weights

are estimated we can calculate the return of the portfolio for the next period and this process will continue till the returns are calculated for all periods.

2.2.2. Mean-variance optimization, an out-of-sample approach with constraint on weights ($\widehat{MV}_c$)

This is the same as $\widehat{MV}$ explained above except for that, with this method we seek to limit weights to positive numbers or in the other word, no short-sell is allowed on assets.

2.2.3. Mean variance optimization, a semi out-of-sample approach ($\widehat{MV}$)

In order to purify the effect of mean estimation on the performance of portfolio, we can use known covariance matrix instead of estimating it out-of-sample hence eliminating the effect of probable covariance matrix estimation error. In this approach the mean is estimated from the sample data but the covariance matrix is the known covariance matrix.

2.2.4. Mean variance optimization, a semi out-of-sample approach with constraint on weights ($\widehat{MV}_c$)

This is the same as $\widehat{MV}$ explained above except for that, with this method we seek to limit weights to positive numbers or in the other word, no short-sell is allowed on assets.

2.2.5. Volatility timing, an out-of-sample approach ($\widehat{MVol}$)

This is a version of the class of strategies introduced by Kirby and Ostdiek (2012) in which the weight of each asset is the inverse of its standard deviation and the weights are summed to one, as in the formula:

$$\frac{\sigma_{i,t}^{-1}}{\sum_{j=1}^N \sigma_{j,t}^{-1}} \quad (3)$$

2.2.6. Volatility timing, a semi out-of-sample approach ($\widehat{MVol}$)

In order to purify the effect of mean estimation on the performance of portfolio, we can use known standard deviation instead of estimating it out-of-sample hence eliminating the effect of probable standard deviation estimation error. In this approach the mean is estimated from the sample data but the standard deviation is the known standard deviation.

2.2.7. Variance timing, an out-of-sample approach ($\widehat{MVar}$)

This is very similar to volatility timing, but in this approach the variance is used instead of standard deviation to determine the weights.

2.2.8. Variance timing, a semi out-of-sample approach ($\widehat{MVar}$)

In order to purify the effect of mean estimation on the performance of portfolio, we can use known variance instead of estimating it out-of-sample hence eliminating the effect of probable variance estimation error. In this approach the mean is estimated from the sample data but the variance is the known variance.

3 Empirical Results

Table 1 - Sharpe ratios for empirical data

Monthly	Historical	Multivariate	Multi_threshold	univariate-combination	univariate-threshold-combination	univariate-threshold-combination posForecast-combination
$\widehat{MV}$	0.3019	0.0197	0.0481	0.1330	0.0564	0.34883*
p-value	0.7040	1.0000	1.0000	0.9999	0.9999	0.1694
$\widehat{MV}$	0.32757*	0.0505	0.0756	0.1794	0.2657	0.35146*
p-value	0.3455	1.0000	0.9998	0.9987	0.8518	0.1356
$\widehat{MV}_c$	0.2988	0.2800	0.2461	0.3769*	0.32972*	0.33707*
p-value	0.7506	0.7744	0.9397	0.2148	0.3658	0.2625
$\widehat{MV}_c$	0.33*	0.2441	0.2341	0.38031*	0.32974*	0.34519*
p-value	0.2982	0.9471	0.9664	0.1808	0.3567	0.1629
$\widehat{MV}_{ol}$	0.2375	-0.0532	0.0771	0.1262	-0.0481	0.2424
p-value	0.9984	1.0000	0.9996	1.0000	1.0000	0.9792
$\widehat{MV}_{ol}$	0.2494	-0.0063	0.0083	0.0749	0.0629	0.2494
p-value	0.9952	1.0000	1.0000	1.0000	0.9999	0.9658
$\widehat{MVar}$	0.2203	-0.0119	0.0190	-0.0428	-0.0223	0.2376
p-value	0.9991	1.0000	1.0000	1.0000	1.0000	0.9886
$\widehat{MVar}$	0.2376	-0.0046	0.0948	0.0877	-0.0489	0.2569
p-value	0.9948	1.0000	0.9995	1.0000	1.0000	0.9586

This table shows the sharpe ratio of different optimization methods and different approaches of factor predictability and their p-values compared to 1/N strategy. If Sharpe of each method is more than historical Sharpe ratio, it is bolded and if Sharpe is more than Sharpe ratio of 1/N strategy, it has a star. Sharpe ration of 1/N strategy is 0.3186.

Table 2 - Sharpe ratios for empirical data

Annual	Historical	Multivariate	Multi_threshold	univariate-combination	univariate-threshold-combination	univariate-threshold-combination posForecast-combination
$\widehat{MV}$	0.3044	0.0782	0.0091	0.44057*	0.42763*	0.43426*
p-value	0.6524	0.9998	1.0000	0.0375	0.0249	0.0004
$\widehat{MV}$	0.33268*	0.0098	-0.0196	0.36934*	0.47368*	0.45661*
p-value	0.2430	1.0000	1.0000	0.3573	0.0018	0.0000
$\widehat{MV}_c$	0.3016	0.2879	0.35006*	0.50639*	0.45649*	0.41492*
p-value	0.6945	0.7442	0.2186	0.0000	0.0001	0.0014
$\widehat{MV}_c$	0.33601*	0.32238*	0.3893*	0.49207*	0.4683*	0.4382*
p-value	0.1878	0.4493	0.0382	0.0000	0.0000	0.0000
$\widehat{MV}_{ol}$	0.2361	0.0186	0.0901	0.0262	0.2004	0.36369*
p-value	0.9981	1.0000	0.9994	1.0000	0.9654	0.1005
$\widehat{MV}_{ol}$	0.2481	-0.0309	0.0399	0.0962	0.0035	0.37534*
p-value	0.9943	1.0000	1.0000	1.0000	1.0000	0.0544
$\widehat{MVar}$	0.2220	0.1032	-0.0337	0.0668	0.2453	0.36022*
p-value	0.9985	0.9991	1.0000	1.0000	0.8601	0.1166
$\widehat{MVar}$	0.2393	0.0012	0.0476	-0.0124	0.33398*	0.37907*
p-value	0.9920	1.0000	1.0000	1.0000	0.3887	0.0440

This table shows the sharpe ratio of different optimization methods and different approaches of factor predictability and their p-values compared to 1/N strategy. If Sharpe of each method is more than historical Sharpe ratio, it is bolded and if Sharpe is more than Sharpe ratio of 1/N strategy, it has a star. Sharpe ration of 1/N strategy is 0.3167.

Table 3 - Sharpe ratios for empirical data

2Year	Historical	Multivariate	Multi_thresho ld	univariate- combination	univariate- threshold- combination	univariate- threshold- posForecast- combination
$\widehat{MV}$	0.3011	0.0974	0.2182	0.46481*	0.45566*	0.41052*
p-value	0.7198	0.9996	0.9402	0.0049	0.0011	0.0059
$\widehat{MV}$	0.33261*	0.1145	0.0564	-0.0277	0.48976*	0.44843*
p-value	0.2792	0.9990	0.9999	1.0000	0.0001	0.0000
$\widehat{MV}_c$	0.2990	0.35798*	0.38717*	0.4654*	0.44109*	0.39914*
p-value	0.7501	0.1703	0.0507	0.0005	0.0007	0.0111
$\widehat{MV}_c$	0.33573*	0.38837*	0.41355*	0.47492*	0.4786*	0.4358*
p-value	0.2285	0.0363	0.0083	0.0000	0.0000	0.0000
$\widehat{MV}_{ol}$	0.2354	0.0643	0.0173	-0.0008	0.30847*	0.36167*
p-value	0.9984	0.9998	1.0000	1.0000	0.5851	0.1148
$\widehat{MV}_{ol}$	0.2475	0.0773	0.0567	-0.0253	0.33117*	0.37135*
p-value	0.9953	0.9995	0.9999	1.0000	0.4096	0.0702
$\widehat{MVar}$	0.2229	-0.0544	0.0568	0.0540	0.35508*	0.3551*
p-value	0.9984	1.0000	0.9999	1.0000	0.2324	0.1598
$\widehat{MVar}$	0.2393	-0.0700	0.0502	-0.0384	0.36038*	0.36684*
p-value	0.9921	1.0000	0.9999	1.0000	0.2135	0.0977

This table shows the sharpe ratio of different optimization methods and different approaches of factor predictability and their p-values compared to 1/N strategy.

If Sharpe of each method is more than historical Sharpe ratio, it is bolded and if Sharpe is more than Sharpe ratio of 1/N strategy, it has a star. Sharpe ration of 1/N strategy is 0.3194.

As is shown in table 1 most of the Sharpe ratios are worse than 1/N strategy in all forecasting method except for univariate-threshold-posForecast-combination which is better than both the 1/N and historical methods in all approaches of optimization. This method of forecasting is best helping the investment performance. Generally univariate methods have better performance than multivariate results, but the increase in Sharpe ratios in univariate-combination and univariate-threshold-combination are marginal, so that only 2 optimization approach can result in a higher than 1/N Sharpe ratio.

However when we widen the time-frame of the predictability (table 2 & 3) the results become stronger in confirming the effectiveness of predictability, specially in univariate methods, on investment performance of the optimized portfolio. So that in univariate-threshold-combination and univariate-threshold-posForecast-combination methods, all optimizations approaches result in much better Sharpe ratios than both 1/N and historical methods.

References

- 1) Campbell J. Y. & Thompson S. B. (2007). Predicting Excess Stock Returns Out of Sample: Can Anything Beat the Historical Average?. *The Review of Financial Studies* 21(4), 1509–1531. doi: 10.1093/rfs/hhm055.
- 2) Carhart M. M. (1997). On Persistence in Mutual Fund Performance. *The Journal of Finance* 52(1), 57-82. doi: 10.2307/2329556.
- 3) DeMiguel, V., Garlappi, L., & Uppal, R. (2007). Optimal versus naive diversification: How inefficient is the 1/N portfolio strategy?. *The review of Financial studies*, 22(5), 1915-1953.
- 4) Fama E.F. & French K.R. (2015). A five-factor asset pricing model. *Journal of Financial Economics* 116(1), 1-22. doi: 10.1016/j.jfineco.2014.10.010.
- 5) Kirby, C., & Ostdiek, B. (2012). It's all in the timing: simple active portfolio strategies that outperform naive diversification. *Journal of Financial and Quantitative Analysis*, 47(2), 437-467.

Email: e.adman@ase.ui.ac.ir

Email: m.botshekan@ase.ui.ac.ir

Socio-economic Differentiation in Experienced Mortality Modelling and its Pricing Implications

Pintao Lyu

Risk Managment Dept., Nationale-Nederlanden, Rotterdam, The Netherlands

Ahmad Salahnejhad Ghalehjooghi¹

Risk Managment Dept., Nationale-Nederlanden, Rotterdam, The Netherlands

Abstract

In recent years, increasing availability and quality of individual-level data enables the life insurers to render fair and flexible pricing based on the personal socioeconomic attributes. Yet, the current pricing of many insurers is based on setting assumptions for the experience factor of the portfolio-specific mortality rates over general population mortality rates. In this experience framework, it's not straightforward to price a personal attribute that is available in the portfolio but not available in the general population. To fill in the blanks, our study uses regression models to account for the personal attributes and constructs corresponding differentiated experience factors, which could be easily embedded in current practice. To address the different uncertainty level of the differentiated experienced mortality, we employ the risk-margin pricing and examine how the differentiated mortality can be reflected in the price. We use the salary information as an example of socio-economic attributes and provide the price of a pure endowment contract for a cohort without/with differentiation. We find that the price differentiation is significant for different salary classes. Also as an example in price calculation, for 40 years old male cohort, salary differentiation can result in around 7% discount for the low salary class and 7.9% surcharge for the high salary class.

Keywords: Differentiation, Experience factor, Socio-economic, Experienced mortality, Pricing, Time-consistent, Risk-margin, EIOPA.

Email: p.lyu@uvvt.nl

Email: ahmad.salahnejhad@nn.nl

¹speaker

Examination of reinforcement learning method for hedging

Hirbod Assa¹

Institute for financial and actuarial mathematics

University of Liverpool, UK

Abstract

Machine learning (ML) has become one of the many useful tools that is used in the field of quantitative finance. As this area of research is still young, there are many questions that need to be answered. That essentially include the comparison of the existing techniques with the ones that is suggested by new ML methods. In this talk we are examining the usefulness of the Deep Reinforcement Learning (DRL), comparing it with the delta hedging as a benchmark. Specifically, we use Deep Deterministic Policy Gradient (DDPG) method and see how this method can converge to a correct solution in a Black-Scholes framework.

Keywords: Hedging, Reinforcement learning, Delta hedging

Email: assa.hirbod@gmail.com

¹speaker

Applying Islamic Repo in monetary policy

Amir Hamooni

Mohsen Yavari¹

FaraBourse Iran Co., Tehran, Iran

Abstract

Short-term interest rate can be changed by using Repo whereby calculating duration and convexity of debt security in market and required price change to achieve target rate. r is a change in short-term interest rate.

$$\Delta r = r_{target} - r_{current}$$

Required price change for modification of YTM calculate by securities duration and calculated duration is deferent based on Repo type.

$$-D_k \Delta r k + 0.5 \Delta r k^2 C k = \Delta p k \Delta p k$$

Duration formula for Repo seller who buy security with Repo proceeds is as following:

$$D_{position} = D_{bonds} - D_{repo}$$

Islamic Repo consist of a call option that hand on by buyer to the seller and a put option, which transfer by seller to the buyer. Option in contract will change calculation related to vulnerability position of each side towards interest rate change. Exercising option by central bank is not based on profit, arising from exercising option, but it is based on extent of interest rate convergence.

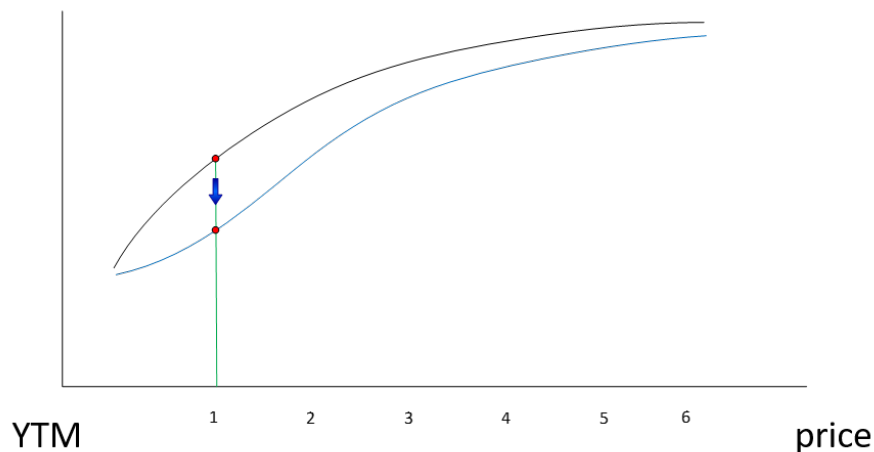


Figure 1

¹speaker

Mapping Capital Requirements to Bank Lending Spreads: The Role of Performance Measures

Ali Golbabaei¹

University of Isfahan, Isfahan, Iran

Mahmoud Botshekan

University of Isfahan, Isfahan, Iran

Abstract

To meet the capital requirements demanded by Basel II and III accords, banks have two main options; namely increasing capital or lowering risk weighted assets. The increasing capital is generally considered as the first option by banks, but as capital is a more expensive source of financing, it causes banks performance measures to deteriorate. One of the main options for a bank to recover its performance measure is to increase lending spreads. King (2010) among others uses return on equity as a performance measure to study this relationship. This study highlights the role of the performance measure in bank decision and extends the model proposed by previous studies to map how increase in the capital ratio impacts on return on equity (ROE) and economic value added (EVA) and how the choice of the performance measures results in different decisions for banks. Using data of Royal Bank of Canada as a case study, the results confirm the conceptual advantages of EVA against ROE as a performance measure and we conclude that EVA can help banks to make better and less expensive decision in the balance sheet management like lending spreads changes. By adding NSFR as liquidity requirement proposed under Basel III to the requirements, the results also confirm the superiority of EVA against ROE. Finally we propose EVA per capital as a more comprehensive measure in banking decisions that also consider valuation considerations for banks.

Keywords: Banks, Basel regulations, Economic value added, Lending spreads

Mathematics Subject Classification [2018]: G21, G28, E51

1 Introduction

The lack of ability in absorbing major on and off-balance sheet risks, as well as deficiencies to consider derivative related exposures during financial crisis resulted in the failures and losses of so many banks. So, in response to these shortcomings, the Basel committee in December 2010 completed a number of critical reforms to the Basel II framework which resulted in new accord named Basel III. The new version introduced new requirements on the regulatory capital, the liquidity risk management and the leverage ratio and banks had to improve their capital ratio and fulfill other requirements up to 2019 [2]. To meet the more restricted capital requirements, banks have two main options: 1- increasing capital 2- lowering risk weighted assets (RWA) [1]. Although increasing capital is generally considered a good deleveraging strategy by regulators, but it results in the bank performance measures to deteriorate. The other option namely shrinking RWA by scaling down loan portfolio (selling assets) or replacing riskier (higher-weighted) loans with safer ones [3] has potentially adverse effects on real economy if many banks simultaneously engage in cutting lending [4]. Cohen and Scatigna (2016) also show that retained earnings (as an internal source to increase capital)

¹speaker

account for the bulk of capital ratio changes, while risk weighted assets reduction plays a lesser role and on average banks tend to increase their lending activities. Therefore banks usually choose to increase capital to meet the Basel requirements.

As capital is a more expensive source of financing, it causes banks performance measures to deteriorate and one of the main options for banks to recover their performance measure is to increase lending spread. King (2010) among others uses return on equity (ROE) as a performance measure to study this relationship, but ROE does not consider the cost of financing that can play an important role in financing decisions. A more comprehensive measure is economic value added (EVA) that also consider the cost of financing, especially if we consider that deleveraging banks can reduce the cost of equity financing. Following the model proposed by king (2010) and using accounting relationship, this study uses a method to map the impact of higher capital and liquidity requirements on bank's lending spreads comparing these two performance measures, namely, ROE and EVA. We show that using different performance measures results in different changes in lending spreads² and using EVA leads to lower lending spread increase. Furthermore we show EVA per Capital can be better performance measure against other traditional measures because it also consider the shareholders' wealth and help banks to make better and less expensive decisions (lending spreads changes) about the balance sheet management.

2 Methodology

In this study, we assume that the bank chooses a targeted capital ratio each year to meet higher requirements required by Basel accords. The total capital ratio (TCR) can be calculated based on the following formula:

$$TCR = \frac{E}{RWA} \quad (1)$$

This formula is divided into two parts: 1- Equity capital (E) 2- The risk weighted assets (RWA). Banks can increase the capital ratio either by rising capital or lowering RWA. As mentioned, we assume that the bank fulfill higher capital ratio by increasing equity and the size and composition of the balance sheet (Risk weighted assets) remain unchanged. In this case the capital ratio changes can addressed as:

$$\Delta TCR_{t+1} = \frac{\Delta E_{t+1}}{RWA_{t+1}} \quad (2)$$

By resolving equation (2) for E_{t+1} , we have:

$$E_{t+1} = E_t + \Delta TCR_{t+1} \times RWA_{t+1} \quad (3)$$

E_{t+1} shows the new equity level needed to meet higher capital ratio. To hold the size of the balance sheet constant, the increase in shareholders' equity can be replaced by decreasing in debt:

$$\Delta equity = -\Delta debt \quad (4)$$

or

$$\Delta TCR_{t+1} \times RWA_{t+1} = -\Delta debt_{t+1} \quad (5)$$

Because debt is substituted with more expensive financing source, namely the equity, increasing the capital ratio by equity, we expect that performance measures like ROE and EVA are reduced. So, the bank needs to raise its landing spreads (among other measures it could take) to recover its performance measures. King (2010) examines the impact of higher capital and liquidity requirement on lending spreads using ROE as a main performance measure for banks. Following model proposed by king (2010), this survey promotes this model and introduces new model based on EVA to show that banks can use EVA as a main performance measure and managers can consider EVA in decision making.

²The lending spread is difference between the interest rate charged on loans and the cost of debts [5].

2.1 Mapping higher capital ratio and EVA to lending spreads

Economic value added is a common measure of shareholder value creation. The EVA is a value created by a firm on its current investment and calculated by difference between net operating profits after taxes and the cost of capital [10]. In fact EVA measures if the managers can generate extra values with respect to cost of financing. The formula for calculating EVA is as following:

$$EVA_{company} = NOPAT - (WACC \times capital) \quad (6)$$

Where $NOPAT$ is net operating profit after tax, $WACC$ is weighted average cost of capital and Capital is the sum of debt and shareholder's equity. Comparing to the other performance measures like return on equity (ROE) and return on assets (ROA), the measure is considered as a more comprehensive measure as it also consider the cost of equity³.

Economic value added is also a tool that banks can use it for measuring their financial performance. But Banks activities are different from activities of other non-financial institutions. One important difference between financial institution and other firms is debt. For non-financial institutions, debt is an important part of financing and therefore in calculation of WACC, Interest expenses must be included. But the liability side of the bank's balance sheet is a part of the business operations of the bank and it is not a financial source [7]. So, the term debt for banks is different with that for firms. In analogy, Interest expenses for banks, on this view is the equivalent of the cost of goods sold for firms. Therefore, in calculation of EVA for banks, it must be excluded in WACC⁴. Consequently, the following economic value added equation (economic profit or equity approach EVA) in equity level is suggested for banks⁵:

$$EVA_{equity \text{ approach}} = net \text{ income} - (C_E \times equity) \quad (7)$$

Consequently, Substituting (3) into (7) yields the complete model:

$$EVA = (NetIntIncome_{t+1} + NonIntIncome_{t+1} - OpExp_{t+1}) \times (1 - tax) - [C_E \times (E_t + \Delta TCR_{t+1} \times RWA_{t+1})] \quad (8)$$

Where $NetIntIncome$ is the net interest income, $NonIntIncome$ is the non-interest income, $OpExp$ is the operating expenses and C_E is the cost of equity. Using the following equation, the net interest income is calculated as:

$$NetIntIncome_{t+1} = [(loan_{t+1} \times a) + (OIntIncome_{t+1})] - [IntExp_{t+1}] \quad (9)$$

Where a is the lending rate, $OIntIncome$ is the other interest income and $IntExp$ is the interest expenses. Using equation (8) and (9) simultaneously, we can back out the lending rate (a) for each targeted EVA. The same line of reasoning can be used to back out the lending spread needed if the ROE used as the performance measure. In this study we assume bank wants to keep previous year performance measure (EVA or ROE) and simultaneously increase its capital ratio. So the question is how the choice of the performance measure can affect the lending rate⁶ required by the bank to prevent its performance to deteriorate. This could be highly important for the bank as in a competitive business conditions, the less lending rate can consider as a competitive advantage for the bank.

³Goldberg (1999) expresses that EVA is different from other traditional measures (like ROE) mainly by taking into consideration the cost of equity. Fiordelisi (2007) states that profit efficiency measures cannot be the most relevant and important performance measure reflecting the bank strategy. Because they don't explicitly take risk or the opportunity cost of capital in consideration. Radic (2015) believes that the term profitability (usually calculated by ROA, ROE) is insufficient to appraise bank stability since it does not consider the level of risk taken by banks. High profitability can be interpreted as signal of bank soundness or high risk taking.

⁴To calculate bank's cost of capital, Maccario et al. (2002) and Franco Fiordelisi (2007), focus on the cost of equity or shareholders' expected rate of return. Since they do not include deposits and other liabilities of the bank in the capital.

⁵Another performance measure is ROE, which is the ratio of net income to shareholder's equity.

⁶In this survey, we just change the lending rate to raise the lending spreads and assume the deposit rate is constant.

3 Main results

In this section we try to answer the question raised in pervious section using the data of Royal bank of Canada during 2016 to 2018. Using data taken from the end of each year (for example the balance sheet, income statement, capital ratio, ROE and EVA), we increase the capital ratio by 1 percentage point for the next year. To prevent previous year performance measures (EVA or ROE) from falling, all else being equal, we raise only lending spreads in response to higher capital ratio. Panel A shows that by the use of information in the end of 2018, 1 percentage point (pp) increase in the capital ratio can reduce ROE from 16 to 15 percent. To recover ROE, the banks needs to raise lending spreads by 10.82 basis point (bp) for the coming year. On the other hand, when we use EVA as a performance measure, the representative bank needs to raise lending spreads by 6.47 bp (Panel B) to recover EVA. We also consider another advantage of EVA, considering varying cost of equity, in response to the higher capital requirements. Panel C, shows that 1 percentage point increase in the capital ratio can reduce the cost of equity from 8.5 to 8.2 percent. In this state, the banks needs to raise only 0.25 bp to recover EVA. The results are also the same in other years.

Table 1: The effect of 1% increase in the capital ratio on lending spreads

Panel A: ROE				
	ROE after 1% increase (%)	Lending spread (BP)	ROE after lending spreads changes	cost of equity (%)
2018	15	10.82	16	
2017	15	19.73	16	
2016	14.7	10.70	15.4	
Panel B: EVA (Fixed cost of equity)				
	EVA after 1% increase (million dollar)	Lending spread (BP)	EVA after lending spreads changes	cost of equity (%)
2018	5595	6.47	5881	8.5
2017	4940	7.74	5262	8.5
2016	4035	6.49	4358	9
Panel C: EVA (Varying cost of equity)				
	EVA after 1% increase (million dollar)	Lending spread (BP)	EVA after lending spreads changes	cost of equity (%)
2018	5870	0.25	5881	8.2
2017	5229	0.78	5262	8.1
2016	4331	0.54	4358	8.6

The results are also robust when we add NSFR as liquidity requirement and EVA per capital as a more comprehensive measure in banking decisions. Because of restriction in space, the results will be presented upon request.

References

- [1] A.R. Admati, et al, *Fallacies, irrelevant facts, and myths in the discussion of capital regulation: Why bank equity is not expensive*. Vol. 86. Max Planck Inst. for Research on Collective Goods, 2010.
- [2] Board, Financial Stability, Basel Committee on Banking Supervision, *Assessing the Macroeconomic Impact of the Transition to Stronger Capital and Liquidity Requirements*, 2010.
- [3] B.H. Cohen and M. Scatigna, *Banks and capital requirements: channels of adjustment*, Journal of Banking & Finance 69 (2016): S56-S69.
- [4] S.G. Hanson, A. K. Kashyap and J. C. Stein, *A macroprudential approach to financial regulation*, Journal of economic Perspectives 25.1 (2011), 3-28.
- [5] M.R. King, *Mapping capital and liquidity requirements to bank lending spreads*, 2010.
- [6] N. Radić, *Shareholder value creation in Japanese banking*, Journal of Banking & Finance, 52 (2015) 199-207.
- [7] A. Thampy and R. Baheti, *Economic value added in banks and development financial institutions*, IIM Bangalore Research Paper 149 (2000).
- [8] D.G. Uyemura, C. C. Kantor and J. M. Pettit, *EVA for banks: Value creation, risk management, and profitability measurement*, Journal of applied corporate finance 9.2 (1996) 94-109.
- [9] F. Fiordelisi, *Shareholder value efficiency in European banking*, Journal of Banking & Finance 31.7 (2007) 2151-2171.
- [10] F. Fiordelisi and P. Molyneux, *The determinants of shareholder value in European banking*, Journal of Banking & Finance 34.6 (2010) 1189-1200.
- [11] A. Maccario, A. Sironi and C. Zazzara, *Is banks' cost of equity capital different across countries? Evidence from the G10 countries major banks*, Evidence from the G10 Countries Major Banks. SDA Bocconi Research Division Working Paper 02-77 (2002).

Email:m.botshekan@ase.ui.ac.ir

Email:aligolbabaei68@ase.ui.ac.ir

An investigation of demand for agricultural commodities in the presence of future market

Nader Karimi¹

Amirkabir University of Technology, Tehran, Iran

Hirbod Assa

Institute for financial and actuarial mathematics, University of Liverpool, Liverpool, UK

Erfan Salavati

Amirkabir University of Technology, Tehran, Iran

Abstract

The purpose of this paper is threefold. First, by assuming that all investors have the possibility of storing goods, we introduced a continuous time speculative storage model that is based on a basic version of a model developed in Deaton and Laroque (1992, 1995, 1996), which incorporates speculative storage into the demand and supply, establishing the concept of stationary rational expectations equilibrium (SREE). Second, we prove that the speculators expectation price obtained from an ODE solution and converges to the market price (SREE). At the end, we estimate the model parameters values in the both without storage case and with storage case, and we show that introduced model is significant by doing (Log) likelihood ratio test. The model proposed in this paper outperforms the data.

Keywords: Stationary Rational Expectations Equilibrium (SREE), Speculative Storage, Likelihood Ratio Test.

1 Introduction

Commodities are a very different asset class from the more traditional classes of traded assets such as equities and bonds. Commodities normally encompass physical goods such as oil, gas, electricity, metals, agriculturals and live stock. The physical nature of commodities is perhaps one of their most defining characteristics specifically because it plays an important role in the behavior of their prices in both the spot and future markets. It is notable that trading with, as well as, presence in the future markets, cause that both demander and supplier immune themselves against the risk of market variations. On the other hand such strategy enables the commodity suppliers to speculate that demands throughout these markets, and take appropriate action in due course. Accordingly, one strategy to be adapted is to have a good perception and knowledge of the commodities demands and price variations, as well. Such strategy enables the commodities

¹speaker

owners to decide, based upon the commodities demands, whether to sell their commodities at the present time or in the future.

The storage models for commodity prices date back to [5], and were further developed by incorporating rational expectations in [6] and [10]. [2, 3, 4] developed a partial equilibrium structural model of commodity price determination and applied numerical methods to test and estimate the model parameters, confronting for the first time the storage model with the documented behavior of actual prices. More recently, many authors improved the storage model in order to better capture the statistical characteristics; see for instance [12, 8, 11, 9, 15, 13] and [1].

2 Model

The aim of this section is to introduce the model base on a rational decision maker problem in the presence of the future market. Let us assume the commodity is traded in a market where the speculation is possible by transferring the purchasing of the good to the next period within some future contracts. Therefore, the spot prices will be impacted by future prices in the following manner: either the producer sells the good at market prices since it is high enough, or not and the producer will account for future market prices. Denoting the market spot prices by p_t and the demand by x_t , one can consider the following rational decision making dynamic ruling the market:

$$p_t = \max\{P(x_t), e^{-r}\mathbb{E}(p_{t+1})\}. \quad (1)$$

Where $\mathbb{E}_t$ denotes the expectation given the information at time t , r is the discount factor, and $P(x_t)$ is the (inverse) demand function. Assa[1], he has introduced the continues time dynamics under the market condition as follow:

$$\begin{aligned} I_t &= x_t - P^{-1}(p_t), (Budget\ Constraint) \\ x_{t+1} &= (1 - \delta)I_t + z_{t+1}, (Equilibrium : Demand = Supply) \end{aligned} \quad (2)$$

where I_t , z_t , x_t , δ denote the inventory, harvest shock, amount-in-hand, depreciation rate and p_t as (1), respectively.

As one can see in (2) the harvest shock coefficient is fixed. While this seems to be a weak assumption, so we addition the scale term to the market condition as follow:

$$\begin{aligned} I_t &= x_t - P^{-1}(p_t), (Budget\ Constraint) \\ x_{t+1} &= (1 - \delta)I_t + x_t z_{t+1}, (Equilibrium : Demand = Supply) \end{aligned} \quad (3)$$

Now, we use of the method which is introdused by Assa [1] to obtain a continuous version of the demand process. If we combine the budget constraint with the equilibrium relation, we get the following

$$x_{t+1} = (1 - \delta)(x_t - P^{-1}(p_t)) + x_t z_{t+1}, \quad (4)$$

we can rewrite (4) as

$$x_{t+1} - x_t = -\delta x_t - (1 - \delta)P^{-1}(p_t) + x_t z_{t+1}, \quad (5)$$

Assa[1], he has assumed that the demand is not affected by changing the framework from discrete time to continuous time. The reason is that the demand function is given on a yearly basis, and not in continuous time. So, with this assump, we have

$$dx_t = (-\delta x_t - (1 - \delta)P^{-1}(p_t)) dt + x_t d(mt + \sigma B_t). \quad (6)$$

where, the parameter m represents the equilibrium or mean value supported by fundamentals, σ the degree of volatility around it caused by shocks, and δ the rate by which these shocks dissipate and the variable reverts towards the mean.

If we denote the discounted price process by $\tilde{p}_t = e^{-rt}p_t$ then we simply have

$$\tilde{p}_t = \max\{e^{-rt}P(x_t), \mathbb{E}(\tilde{p}_{t+1})\}. \quad (7)$$

Now we want to find a solution for this equation. In principle, there is no closed form solution for this price dynamic however, with some changes and making some assumptions one can get a fair approximation of the solution.

First of all consider the economy has only finite periods, say T . In that case the commodity producer needs to sell everything at the market price as T . As a result, we have $\tilde{p}_T = e^{-rT}P(x_T)$. From the construction of $\tilde{p}_t$ one can see that it is nothing but the Snell envelop of the process $z_t = e^{-rt}P(x_t)$ with $\tilde{p}_T = e^{-rT}P(x_T)$. In order to study this process we consider then a continuous time version of the model by introducing a continuous Snell envelope as follow

$$\tilde{p}_t = \sup_{t \leq \tau \leq T} E(z_\tau | \mathcal{F}_t) = \sup_{t \leq \tau \leq T} E(e^{-r\tau}P(x_\tau) | \mathcal{F}_t), \quad (8)$$

where the supremum runs over the set of all stopping times. However, we need to make another assumption in our economy, to solve the pricing process: we assume the prices are at the steady state. That means one can assume the price process is independent from time and it is only a function of the state variable $x_0 = x$:

$$\tilde{p}(x) = \sup_{t \leq \tau} E(z_\tau | \mathcal{F}_t) = \sup_{t \leq \tau} E(e^{-r\tau}p(x_\tau) | \mathcal{F}_t), \quad (9)$$

Now let us assume that P is simple linear market (inverse) demand function define as

$$P(x) = (b - ax)_+ \quad (10)$$

for some $a, b > 0$ (Fig. 1).

However, this problem can be solved in an optimal stopping time framework. First of all observe that $\tilde{p}(x) = \sup_{\tau} E(e^{-r\tau}(b - ax_\tau)_+ | \mathcal{F}_t)$ where x_t follows a geometric brownian motion. If the demand process starts at x , i.e., $x_0 = x$, and we denote $V(x) = \sup_{\tau} E(e^{-r\tau}(b - ax_\tau)_+ | \mathcal{F}_t)$ then by Theorem(11.2.1) in [14], it is known that V must solve the following problem

$$\begin{aligned} \mathbb{L}_g V &= rV, & V(x) &> P(x) \\ & & V(x) &\geq (b - ax)_+ \end{aligned} \quad (11)$$

Where $\mathbb{L}_g$ is as follow:

$$\mathbb{L}_g = ((m - \delta)x_t - (1 - \delta)P^{-1}(g(x_t))) \frac{\partial}{\partial x} + \frac{1}{2}\sigma^2 x^2 \frac{\partial^2}{\partial x^2}. \quad (12)$$

This is the demand function of the good price in the presence of the future market.[see. Fig.(2)]

Remark 2.1. When the market get to equilibrium, the $V(x)$ function Should be equal to the market price. i.e. $V(x) = f(x)$. Following (Deaton), such function f a stationary rational expectations equilibrium(SREE).

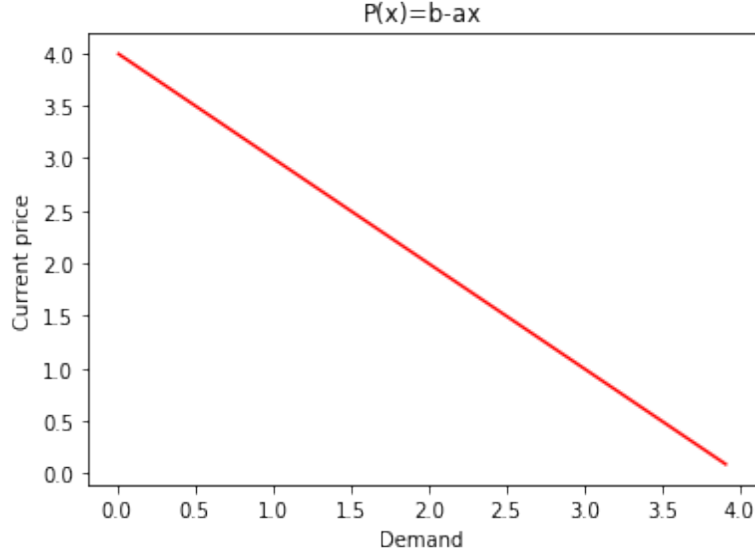


Figure 1: Graph of (inverse) demand function

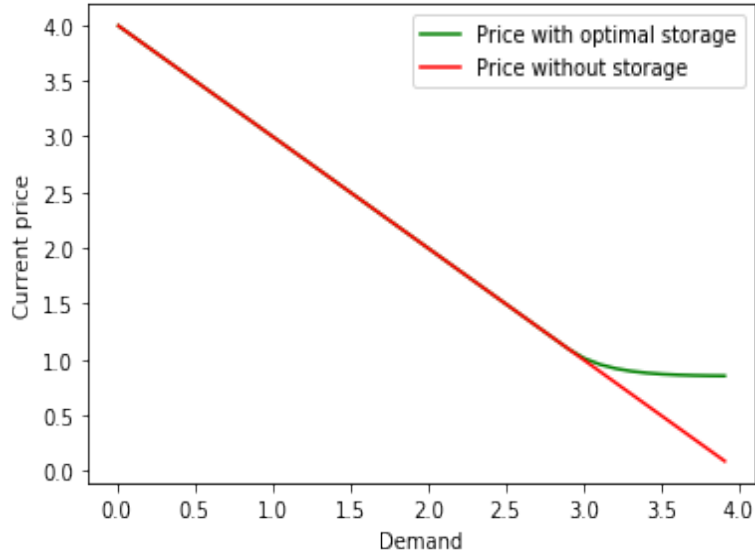


Figure 2: Prices with and without storage

Theorem 2.2. *There is a unique SREE f in the class of non-negative continuous non-increasing functions. Furthermore, let $x^* = \sup\{x : f(x) = P(x)\}$. Then,*

$$\begin{cases} f(x) = b - ax & \text{when } x \leq x^* \\ f(x) > b - ax & \text{when } x > x^* \end{cases} \quad (13)$$

f is strictly decreasing whenever it is strictly positive.

3 Simulation of the theoretical model

In this section, to provide a better understanding of the proposed algorithm, we itemize a simple pseudo code to show how to choose the function value of f and estimated values (m, σ, x^*) . The algorithm runs as follows:

1. Start by a initial values $(f_0 = P, a, b, \delta, m, \sigma, r)$.

2. Solving the ODE (14) on the 400 points with known initial conditions and receive the new V on the every step, the sequence V_n converges to the SREE f .

So, let $f_0 = P$, and $x_n^* = \inf\{x : \mathbb{L}_{V_n} V_n(x) > rV_n(x)\}$, we should solve the following ODE:

$$\begin{cases} \mathbb{L}V_{n+1} = rV_{n+1}, x > x_n^* \\ V_{n+1}(x_n^*) = P(x_n^*) \\ V'_{n+1}(x_n^*) = P'(x_n^*) \end{cases} \quad (14)$$

in which $\mathbb{L}_{V_n}$ satisfies in (12). By utilizing the Algorithm, we have the following descritization:

$$\begin{cases} V_{i+2} = \left(\frac{2K_i \Delta x}{\sigma^2 x_i^2} + 2 \right) V_{i+1} + \left(\frac{2(\Delta x)^2}{\sigma^2 x_i^2} r - \frac{2K_i \Delta x}{\sigma^2 x_i^2} + 1 \right) V_i \\ V_{i^*} = P(x_{i^*}) \\ V_{i^*+1} = P(x_{i^*+1}) \end{cases} \quad (15)$$

Where and $K_i = ((m - \delta)x_i - (1 - \delta)P^{-1}(f(x_i)))$, and $i^* = \inf\{i : \mathbb{L}V_{old}(x_{i^*}) > rV_{old}(x_{i^*})\}$.

4 Estimation parametes with simulated data and actual data

Duo to demand function heavily depend on the demand components of the market i.e. σ, m . But, due to the Therefore, we estimate these parametes in two cases. The first case we estimate those with simulated data , the second with actual data.

A. Estimate with simulated data

We use simulated data in two cases of demand functions (Linear and Isoelastic) with some initial values. Table 5 and Table 6 show the (Log) likelihood value and parameters estimates for various models (with and without storage), respectively .

Table 5
Log Likelihood Values and Parameters Estimates For Linear Demand.

	Without-Storage	With-Storage
m	4.326	6.97
σ	0.5608	0.98
x^*	-	2.87
MLE	-8.22647e+08	-1.438844e+04

Table 6
Log Likelihood Values and Parameters Estimates For Isoelastic Demand.

	Without-Storage	With-Storage
m	4.48	6.68
σ	1.64	2.013
x^*	-	16.82
MLE	-7.1352e+0.6	-2.438844e+03

B. Estimate with actual data

The actual data used in this study come from the actual commodity data (Sugar, Rice, Wheat, Tea and Soybeans Monthly Price - US Dollars per Kilogram from Oct 1917 to Oct 2018 deflated by the U.S). consumer price index. We show the estimated parameters and the likelihood function value in Table 7 and Table 8, for Linear demand and Table 9 and Table 10 for Isoelastic demand, respectivity.

Table 7
Parameters Estimates with Actual Data For Linear Demand.

	m	$m_{Storage}$	σ	$\sigma_{Storage}$	x^*	$x_{Storage}^*$
Sugar	2.89	3.64	1.41	2.01	-	2.85
Rice	3.42	4.79	1.05	1.2	-	2.31
Tea	6.35	7.88	1.24	1.91	-	2.73
Wheat	5.01	5.94	1.64	1.99	-	2.51
Soybeans	3.71	4.98	1.53	1.98	-	2.67

Table 8
Log Likelihood Values For Linear Demand.

	Without-Storage	With-Storage
Sugar	-3.2138e+06	-2.6657e+02
Rice	-2.1524e+04	1.2234e+02
Tea	-5.7418e+05	-2.6529e+03
Wheat	-8.7102e+06	-1.5264e+02
Soybeans	-3.9773e+04	-1.03149e+04

Table 9
Parameters Estimates with Actual Data For Isoelastic Demand.

	m	$m_{Storage}$	σ	$\sigma_{Storage}$	x^*	$x_{Storage}^*$
Sugar	3.33	3.89	1.02	1.95	-	1.3
Rice	3.12	4.86	1.001	1.16	-	1.24
Tea	6.87	7.91	1.11	1.88	-	2.1
Wheat	4.96	5.79	1.23	1.94	-	2.87
Soybeans	3.66	4.78	1.68	2.001	-	2.5

Table 10
Log Likelihood Values For Isoelastic Demand.

	Without-Storage	With-Storage
Sugar	-2.15634e+03	-1.652824e+01
Rice	-3.23181e+04	-1.141852e+02
Tea	-7.54618e+05	-2.433696e+03
Wheat	-9.24732e+04	-1.5216e+04
Soybeans	-8.22647e+08	-1.43884e+04

References

- [1] Assa H, (2015) *A financial engineering approach to pricing agricultural insurances*, Agricultural Finance Review, Vol. 75 Issue: 1, pp.63–76.
- [2] Deaton, A. and G. Laroque (1992). On the behaviour of commodity prices. The Review of Economic Studies 59(1), pp. 1–23.
- [3] Deaton, A. and G. Laroque (1995, Suppl. De). Estimating a nonlinear rational expectations commodity price model with unobservable state variables. Journal of Applied Econometrics 10(5), S9–40.
- [4] Deaton, A. and G. Laroque (1996, October). Competitive storage and commodity price dynamics. Journal of Political Economy 104(5), 896–923.

- [5] Gustafson, R. L.(1958). Implications of recent research on optimal storage rules. *Journal of Farm Economics* 40(2), pp. 290–300.
- [6] Muth, J. F.(1961). Rational expectations and the theory of price movements.*Econometrica* 29(3), pp. 315–335.
- [7] Samuelson, P. A. (1971). Stochastic speculative price. *Proceedings of the National Academy of Sciences of the United States of America* 68(2), pp. 335–337.
- [8] Ng, S. and F. J. Ruge-Murcia (2000). Explaining the persistence of commodity prices. *Computational Economics* 16, 149–171.
- [9] Miao, Y. W. W. and N. Funke (2011). Reviving the competitive storage model: A holistic approach to food commodity prices. *IMF Working Paper*.
- [10] Samuelson, P. A. (1971). Stochastic speculative price. *Proceedings of the National Academy of Sciences of the United States of America* 68(2), pp. 335–337.
- [11] Cafiero, C., E. S. Bobenrieth, J. R. Bobenrieth, and B. D. Wright (2011). The empirical relevance of the competitive storage model. *Journal of Econometrics* 162(1), 44 –54.
- [12] Chambers, M. J. and R. E. Bailey (1996). A theory of commodity price fluctuations. *Journal of Political Economy* 104(5), pp. 924–957.
- [13] Arseneau, D. and S. Leduc (2012). Commodity price movements in a general equilibrium model of storage. preprint.
- [14] B. Oksendal, *Stochastic Differential Equations: An Introduction with Applications*, Springer Science and Business Media, 2003.
- [15] Evans, L. and G. Guthrie (2007). Commodity price behavior with storage frictions. pp. 1451 – 1468.

e-mail: nkarimi@aut.ac.ir

e-mail: H.Assa@liverpool.ac.uk

e-mail: erfan.salavati@aut.ac.ir

An Introduction to Theoretical and Historical Aspects of Evolutionary Economics and Evolutionary Approaches in Economic Modeling

Ebrahim Kargar

PhD in Economics, CEO of Dana Insurance

Mehdi Ahrari

PhD student in Economics, Allameh Tabataba'i University, Director of Dana Insurance Research and Development Bureau

Abdollah Najar Firozjahi¹

PhD in Economics Student of Tarbiat Modares University, Expert of Dana Insurance Research and Development Bureau

Abstract

The fundamental problem that presumably expanded alternative streams and schools in economics in the course of reform and evolution of mainstream and orthodox economics was the consideration of facts and roots of micro and macroeconomic phenomena, which are entirely influenced by humankind as its central axis, and significantly affect social and economic environments as well as associated institutions and institutional structures. However, the major concern in most studies, which is still poorly understood, is to reflect the epistemological and methodological bases of modeling in the fields of humanities (in general) and economics and insurance (in particular) that realistically constitute the set of studies conducted conventionally and carries the technical innovations in the modeling within the same conventional streams. In other words, the distinct aspects that arose from the application of new methods to modeling are entirely based on the method substitution, and no to deepen the cognitive foundations in the fields of epistemology and methodology, which have practically created no evolution in this course. This paper attempts to provide knowledgeable and profound insight into the modeling by describing the evolutionary literature in economics and explaining its theoretical aspects. The GMDH algorithm is ultimately presented as an example of evolutionary approaches with a well-established background in economic studies

Keywords: evolutionary economics, evolution, evolutionary methods, modeling, GMDH algorithm

1 Introduction

ing a valid scientific apparent and causal model. Realistic scientists conclude that science is not aimed at prediction, because the admission of this goal is equal to the concentration of technology and setting cognition aside and restricting to the instrumental aspects of scientific. The analysis and evaluation of any human consequence entail a precise and comprehensive understanding of the inherent characteristics and dimensions as well as endogenous and exogenous factors, although it does not make all the cognitive aspects clear to us. A principal problem while analyzing human phenomena is to recognize the roots of the subject under investigation. Such recognition must have two central components, including comprehensiveness and methodology. Since human phenomena are multidimensional with several aspects of various kinds, its cognitive analysis must, therefore, be based on a kind of subjective and inter-discipline comprehensiveness. The second component, which is indeed a key to solve the problem, is adopting an optimal methodology

¹speaker

that is well suited to the subject under investigation. Such a methodology manages problem analysis and explanation, besides providing a way to find the cognitive origins of the subject. These two components describe a kind of developmental and institutional approach that has been cited by many economic researchers in recent decades as a noticeable way of identifying and explaining human phenomena (in general) and economic phenomena (in particular). A prominent feature of the methodology of heterodox thought streams, and particularly institutionalists, is the importance of cognitive studies while analyzing and explaining human phenomena. The cornerstone of the new institutional economics includes the conventional neoclassical theory (i.e. methodological individualism) and the principle of utilitarianism. Any change in the methodological foundations of an economist will affect his/her understanding of the facts and his conceptualization of economic phenomena, and, strictly speaking, the foundations adopted in the field of the methodology of interest is the origin for the plurality of methodologies and, hence, economic perceptions. On the other hand, there is a need for a balance between ontological assumptions and methodological principles to understand socio-economic realities so that, ontologically, the social reality is considered multi-layered and evolving, and to avoid explaining the subject by a single layer. Theories must also cover the evolution of realities and not be permanently extracted by establish research. Therefore, numerous fundamental principles need to be analyzed to develop a precise methodology. The first essential principle is realism and the avoidance of abstract and subjective presuppositions, frequently without any dimension of the subject facts. The second is the consideration of the plurality and diversity of the phenomenon of interest, which necessitates the adoption of interdisciplinary approaches. Avoiding limited and dualistic attitudes and concentrating on the interactive relationship perceptions is the third principle resulting from the methodology of human subjects (in general) and economic problems (in particular). The fourth fundamental principle of the methodology is to prevent classifying into the framework of a particular theory. We have to think beyond theoretical boundaries and limiting paradigms. This does not mean to excommunicate the valuable theories of scholars and philosophers which all of them have been a light of path for their posterities and means that the recognition is not a limited and static category, but is expanded, dynamic, and profound which constantly calling humankind. The following is an overview of the evolutionary economics as well as behavioral and laboratory economics that recently awarded the Nobel Prize. The evolutionary algorithm GMDH is ultimately discussed as a prosperous method in the field of economic studies.

2 Evolutionary Economics

Evolutionary economists believe that the dynamic economy is constantly changing and in turmoil, not always tending towards equilibrium. Creating and providing resources for those commodities involves many processes that change with the advance of technology. Organizations that govern these processes and production systems, as well as consumer behavior, must evolve as the process of production and procurement changes.

3 Laboratory Economics

Laboratory economics is nowadays an accepted method for testing hypotheses and validating economic models. The core of laboratory economics is game theory, and laboratory economics is an empirical framework for running games and comparing scientific results with theoretical predictions. Laboratory economics is used in a wide range of topics, including market mechanisms, evolutionary theory, the decision-making process, financial economics, learning and many economic and social issues

4 behavioral economy

Behavioral economics is the knowledge that attempts to better describe and analyze our economic behaviors and decisions by combining our economic knowledge and the gains of psychology knowledge, and in particular cognitive psychology. The behavioral economics of the first relied heavily on evidence formed by experiments, but recently behavioral economists have gone beyond experimentation and used the methods of economists

(Camerre, Lunstein Mathew, 2004). Recently, field data has greatly helped behavioral economics. Many articles have used techniques such as field tests, computer simulations, and even brain scans.

5 GMDH Algorithm

The consideration of technical and quantitative fields should concentrate on developing diverse, competing, and collaborative perspectives to design, simulate, and extract patterns that are close to the realities of system dynamics, besides analyzing and recognizing economic phenomena in terms of their behavioral nature and inherent dynamics, all to predict similar future consequences. Evolutionary algorithms, neural networks, and other intelligent methods have been surely able to propose integrated patterns (i.e. the application of regression results as neural network inputs), which provide high explanatory and predictive capacity, besides fixing some of the inherent limitations of conventional econometric methods. From a methodological perspective, the consideration of ontological aspects and stress on process explanation, following the existential nature of the phenomenon under consideration and the principle of nonlinear dynamics, is a prominent feature of the foregoing methods. Most of the phenomena we are practically encountering are nonlinear in nature and the equation governing their performance is difficult to obtain. What we have in cash is a large quantity of time-series data. Accordingly, the application of dynamic nonlinear systems in analyzing economic time-series is of great importance in the more documented and realistic analysis and predictions. Dynamic nonlinear systems exhibit a variety of behaviors that can be applied to justify many of the seemingly random economic phenomena. As such, many attempts have been made to fit the time-series of nonlinear dynamic functions

References

- [1] Azad (Armaki), Gholamreza, 2005, Psychology of Economics, Tehran, Ney Publication.
- [2] Ahrari, Mehdi, Motaseli, Mahmood (2015), Explaining Life Insurance from a Methodological Perspective, Journal of Comparative Economics, Humanities and Cultural Studies Institute, Volume 2, Issue 1, Spring and Summer 2015.
- [3] Tohidfam, Mohammad, 2011, Rethinking the Foundations of Evolutionary Hayek Epistemology, Journal of Political and International Approaches, Volume 1, Number 25.
- [4] Judge, Yadollah, 2007, Introduction to Methodology of Economics, Tehran, Ney Publication.
- [5] Sanjani, Karim, 2002, History of Economic Beliefs, Tehran, Tehran University Press.
- [6] Jalili Marand et al., 2016, Modeling Social Preferences in Laboratory Economics: Laboratory Introduction and Investigation, Economic Research.
- [7] Nazarpour, Mohammad Naghi et al., (2015), Application of Laboratory Economics in Bank Resource Allocation (Case Study of Banking Contracts), Journal of Islamic Economics.
- [8] Judge Nouri, Seyyed Sepehr, Narimani et al. (2014), Analyzing the Rational Economics of Conventional Economics in the Field of Policy-Based Science, Technology and Innovation, Journal of Innovation Management.
- [9] Rudolf Richter, Translator: Jalili, Mir Hossein (2008), The New Institutional Economics, Beginning, Meaning and Prospects, Economic Journal of Review and Economic Policies.
- [10] Suzannchi Kashani, Ebrahim, Safae, Navid, 1396, Introduction to the epistemology of evolutionary economics; Investigating the feasibility of the applicability of biological evolution concepts in the socio-economic sphere, Humanities Methodology, Volume 23, Issue 91
- [11] Andreoni, J., Bernheim, B.D. (2009). Social and the 50-50 Norm: *A Theoretical and Experimental Analysis of Audience Effects*. *Econometrica*, 77(5), 1607-1636.

- [12] Cappelen, A. W, Hole, A. D., Sorensen, E. Q., Tungodden, B. (2007). The Pluralism of Fairness Hdeals
- [13] Korsi, Dan; Experimental Economics; Working paper: Center for Public Policy Studies, School of Public Policy, University of Chicago, 2009.
- [14] Reuben, Ernesto; Experimental Economics; Provided in experimental economics at Columbia Business School press, 2013.
- [15] John H. Kagel and Alvin E. Roth; The Handbook of Experimental Economics; USA: editors, Princeton University Press, 1995 1997.
- [16] SCOTT, D.E and HUTCHINSON, C.E. (1976). The GMDH Algorithm- A Technique for Economic Modeling. Report No.ECE-SY-67-1, University of Massachusetts, Dep. of Computer Sciece.

Email: omid.najar@yahoo.com

A Survey on Data Mining Methods for Customer Churn Prediction in Insurance Industry

Mitra Ghanbarzadeh¹

Assistant Professor

Insurance Research Center, Tehran, Iran

Asma Hamzeh

Assistant Professor

Insurance Research Center, Tehran, Iran

Abstract

Prediction and detection of customer churn in the insurance industry is a core research topic in recent years. Insurers need to know the reasons of customer churn to predict churners, which can be realized through the knowledge extracted from data using data mining techniques. Therefore, accurate recognition of factors influencing customer churn and well-known data mining methods can help researches to choose methods and variables. The main purposes of this research are to present the main factors of customer churn and review various data mining techniques when applied to customer churn prediction. Besides, we present the strengths and weakness points for churn prediction methods based on reviewing the published papers in this field.

Keywords: Insurance, Data Mining, Customer Churn, Customer Retention, Customer Relationship Management.

Mathematics Subject Classification [2010]: 97M30, 91C20, 92B20.

1 Introduction

Customer Relationship Management (CRM) in the insurance industry is concerned with the relation between the policyholders (as customers) and an insurance company (or insurer). CRM is a very broad discipline, it reaches from basic contact information to marketing strategies which leads to create superior value for the insurer and the policyholders. CRM can be viewed in four aspects which are: customer identification, customer attraction, customer development and customer retention (Ngai et al., 2009, Huigevoort, 2015).

An example of customer identification is customer segmentation, e.g. based on gender. Customer attraction deals with marketing-related subjects such as direct marketing. An important element of customer development is the up-selling sales technique. Finally, customer retention is the central concern of CRM, and is linked to loyalty programs and complaints management. Customer satisfaction, which refers to the difference in expectations of the customer and the perception of being satisfied, is the key element for retaining customers (Ngai et al., 2009, Huigevoort, 2015). Customer retention is about exceeding customer expectations so that they become loyal to the brand. For retaining customers, we need to study and analyze the main reasons of customer churn. Customer churn is a big issue of the insurance industry. The churn means those customers who want to leave the policy in the near future. As an insurer, every policyholder retained is as successful and an indication of a better future view.

Therefore, there is an essential need to predict those customers on behalf of some parameters to perform some suitable action to minimize their churn.

Customer churn in the insurance industry has divided into three types:

¹ speaker

- Involuntary churn: This occurs when policyholders fail to pay policy premiums and as a result, the insurer terminates the policy.
- Inevitable churn: This occurs when policyholders die or migrate resulting in omitting customers from the market completely.
- Voluntary churn: This occurs when policyholders prefer to switch to another insurer because of more value.

In the meantime, how to predict and prevent the churn of customers, has received the attention of many insurers and researchers in this field.

The main purposes of this research are presentation main factors of customer churn and review various data mining techniques when applied to customer churn prediction. Besides, we compare the attained methods based on their strengths and weakness points and finally, the papers that used data mining to detect customer churn are analyzed.

The paper is organized as follows. Section 2 reviews the main variables that effecting customer churn in the insurance industry. Section 3 presents some data mining techniques which can be used to predict and estimate customer churn. Finally, Section 4 concludes the paper.

2 Variables Impacting the Customer Churn

The main step in the data selection procedure to predict customer churn is to extract the variables which can impact customer churn. Therefore, we investigate papers that study these variables which the most important churning characteristics found in this research, based on Huigevoort (2015), Risselada et al. (2010), Gunther et al. (2014) and Almana et al. (2014), are presented in Table 1. It should be mentioned these variables are completed, duration our research.

Table 1. Some variables that impact on customer churn

Demographic-related variables	Insurer-related variables	Product-related variants
Identification number Age Gender Marital status Network attributes Segment selected by the company Educational level Income Customer satisfaction Employment status Occupation	Number of complaints Reaction on marketing actions Number of declarations Outstanding charges Duration of current insurance contract Number of times subscribed Competitors with superior technology Type of contact (email, call, ...)	Premium Discount Payment method Type of insurance Product usage Brand credibility Switching barrier Interest rate Sum Insured Claim history Duration of current insurance contract The proportion of fee Number of policy

Before performing the data mining procedure, we need to select data based on the main variables of churn cause and run pre-processing methods on these data. Some data set are incomplete, inconsistent and noisy, which need to improve or eliminate their problem before performing any data mining methods. The pre-processing part in data analysis is necessary to get high-quality data which leads to better data mining results. Data pre-processing involves data cleaning, data imputation, data integration, data transformation, and data reduction.

3 Data Mining Techniques to Predict the Customer Churn

There are different data mining techniques that can predict customer churn in insurance study. Generally, data mining divided into two categories: Predictive data mining and Descriptive data mining. In one hand, predictive data mining models are used to express system which can help to predict the performance of various variables. Therefore, producing a model that can estimate by using executive codes, i.e. ranking is the scope of predictive data

mining. On the other hand, descriptive data mining study new data and performance describe the behavioral patterns of variables based on available data set. Then, a comprehensive understanding of the current system by using its hidden patterns and internal relations of a data set is the purpose of descriptive data mining.

In Table 2, we present some well-known data mining methods to predict and detect customer churn. It should be mentioned, these methods are completed in the paper. Besides, the performance of data mining methods in customer churn prediction is studied in this paper based on published articles in this field.

Table 2. Brief Description of Data Mining Methods

Data Mining Methods	Brief Description
Decision Trees	– The most common method used in predicting and evaluating the classification of customer churn problems – Two steps: tree building and tree pruning – Classify and label records – Alter the tree until a leaf node is attained in the customer's dataset – Assign a customer record to churner or non-churner leaf node
Logistic Regression	– The response variable is binary: churn or non-churn – The remaining variables are mostly continuous in nature because of that logistic regression appeared to be the best choice – Calculate the probability of each client churn
Clustering	– Different types: K-Means Algorithm, fuzzy clustering, hierarchical clustering – Partition the customers into clusters and calculate churn rate for each cluster to attain different profiles
Classification	– The databases can be segmented into more homogeneous groups – Then the data of each group can be explored, analyzed and modeled. – Segmentation can be done using variables associated with risk factors, profits or behaviors.
Neural Networks	– Identify complex relationships within the data – Dependent Variable: Exit of the customers from the insurance company – The input layer of Network: factors influencing customer churn
Support Vector Machine	– Used for classification and regression analysis – Used to classify churners and non-churners – Construct hyper-planes in a multidimensional space to separate churn and non-churn customers

References: Ngai et al. (2009), Huigevoort (2015) and Alman et al. (2014)

To detect which churn prediction methods are widely used in the literature, a literature review is performed. The four most used techniques in the literature are logistic regression, decision tree, neural networks and support vector machines (Huigevoort, 2015), which the main results of comparing these techniques are as follows:

- Models based on logistic regression have been attached with good results in the prediction and detection of customer churn. It is based on a supervised learning model, i.e., a data set of past observation is used to see future values of the explanatory and numerical targeted variables.
- The logistic regression and neural network techniques showed the best performance when applied to pre-processed data. Also, when the normality assumption of the data set is not held, logistic regression is less affected.
- The neural network model is suitable with data which have non-linear behavior or noise component. Two disadvantages of this model are the difficult parameter choices and the difficult interpretation due to these complex relationships.
- Decision trees are not suitable for complex and non-linear relationships between the attributes.

4 Conclusion

Market saturation, increasing the competitiveness of the insurance products and the higher attraction cost of new customers than retaining old customers are the main reasons for keeping customers in insurance companies. To attain this goal, it is necessary to identify factors influencing customer retention, which can be used for the

detection of customer churn. There is probably a lot of research available in the field of data mining applications for customer churn detection, but it is still an active field of research to obtain more accurate solutions. In this paper, we present variables impacting customer churn which divided into three categories: demographic, insurer and product-related variables. Besides, a review of data mining techniques used for customer churn prediction in the insurance industry is studied. The most important data mining methods are neural network, decision tree, logistic regression, support vector machine, clustering, and classification. The review shows to find customer churn prediction depends on the objectives of decision-makers.

The results can enable analyzing and predicting future behaviors by considering the dark and unknown dimensions of customer behavior and considering new approaches to preventing customer churn. Based on the results, the following are some suggestions:

- Value-added services, discounts, and promotional activities to satisfy and loyalty customers,
- Customer segmentation and understanding the needs of each group and forecasting future needs,
- Training of insurance companies staff in marketing skills to provide effective services,
- Importance of service quality,
- Strengthen customer complaints handling system to gain customer satisfaction,
- Direct and indirect surveys to identify customer expectations.
- Organizational structure and human factors can affect customer retention. With the right organizational structure, more clients can be retained for longer in the insurance company.

References

- [1] Almana, A. M., Aksoy, M. S., and Alzahrani, R. A survey on data mining techniques in customer churn analysis for telecom industry, *International Journal of Engineering Research and Applications*, vol. 45, pp. 165–171, 2014.
- [2] Gunther, C.C., Tvette, I.F., Aas, K., Sandnes, G.I. and Borgan, R. Modelling and predicting customer churn from an insurance company. *Scandinavian Actuarial Journal*, 2014(1):58, 2014.
- [3] Huigevoort Ch., *Customer churn prediction for an insurance company*, Master Thesis of Eindhoven University of Technology, 2015.
- [4] Ngai, E.W.T., Xiu, L., and Chau, D.C.K. Application of data mining techniques in customer relationship management: A literature review and classification. *Expert Systems with Applications*, 36(2):2592-2602, March 2009.
- [5] Risselada, H., Verhoef, P.C. and Bijmolt, T.H.A. Staying Power of Churn Prediction Models. *Journal of Interactive Marketing*, 24(3):198, 2010.

Email: ghanbarzadeh@irc.ac.ir

Email: hamzeh@irc.ac.ir

Optimization portfolio under uncertain condition

Sadegh Miri¹, Erfan Salavati

Department of Mathematics and Computer Science, Amirkabir University of Technology, Tehran, Iran

Abstract

The uncertain theory offers firstly with Liu² in the year 2007. The variable on uncertain theory is useful for modeling human uncertainty. we want to find the best weight for each environment such have a Return and Volatility, in the portfolio. The return of variable is an uncertain variable, In More three the person thinks the Return of i'th environment are uncertain variables with Expected value μ_i and Volatility σ_i^2 , Then we use an exponential utility function of Return portfolio for find the best weight of environments.

Keywords: environment portfolio, uncertain variable, utility function, uncertain optimization.

AMS Mathematical Subject Classification [2018]: 91G10, 03B30

1 Introduction

Some information and knowledge are usually represented by human language like “about 100km”, “middle age”, and “big size”. How do we understand them? Perhaps some people think that they are subjective probability or they are fuzzy concepts. However, a lot of surveys showed that those imprecise quantities behave neither like randomness nor like fuzziness. This fact provides a motivation to invent another mathematical tool, namely uncertainty theory.

The uncertain variable is a special type of variable such that we relate to every event the probability uncertain (near be .5). The person in exchange stock market usually according to the feeling, buy and sell the stock. In this article, we want to use the Normal uncertain variable for Return of stock then use the exponential utility function for the return of the portfolio. and As a result, find the best weight of the environment for stocks.

1.1 Uncertain Measure

Let Ω be a nonempty set, and collection $\mathfrak{F}$ are σ -algebra on all subset of Ω .

Definition 1.1. Uncertain measure μ is a function from $\mathfrak{F}$ to $[0, 1]$ which for each event A of $\mathfrak{F}$, The $\mu(A)$ is indicates the belief degree that A will occur. In order to ensure that the number $\mu(A)$ has certain mathematical properties, Liu[1] proposed the following four axioms:

¹speaker

²Baoding Liu

Axiom 1. (*Normality Axiom*) $\mu(\mathfrak{F}) = 1$ for the universal set $\mathfrak{F}$.

Axiom 2. (*Monotonicity Axiom*) $\mu(A_1) \leq \mu(A_2)$ whenever $A_1 \subseteq A_2$.

Axiom 3. (*Self-Duality Axiom*) $\mu(A) + \mu(A^c) = 1$ for any event A .

Axiom 4. (*Countable Subadditivity Axiom*) For every countable sequence of events $\{A_i\}$, we have

$$\mu\left(\bigcup_{i=1}^{\infty} A_i\right) \leq \sum_{i=1}^{\infty} \mu(A_i).$$

Example 1.2. Let $\Omega = \{a_1, a_2, a_3\}$. For this case, there are only 8 events.

Define

$$\begin{aligned} \mu(a_1) &= .6 & \mu(a_2) &= .3 & \mu(a_3) &= .4 \\ \mu(a_2, a_3) &= .4 & \mu(a_1, a_3) &= .7 & \mu(a_1, a_2) &= .6 \\ \mu(\emptyset) &= 0 & \mu(\mathfrak{F}) &= 1 \end{aligned}$$

Definition 1.3. Uncertain variable ξ on Uncertain space $(\mathbb{R}, \mathfrak{B}(\mathbb{R}), \mu)$ is normal with expected value μ and Volatility σ^2 if The Distribution Function of ξ be

$$f(x) = \frac{\frac{\pi}{\sqrt{3}\sigma} e^{\frac{\pi(\mu-x)}{\sqrt{3}\sigma}}}{(1 + e^{\frac{\pi(\mu-x)}{\sqrt{3}\sigma}})^2} \quad (1)$$

hence, we have Cumulative Distribution Function of ξ are

$$\Phi(x) = \left(1 + e^{\frac{\pi(\mu-x)}{\sqrt{3}\sigma}}\right)^{-1}. \quad (2)$$

Theorem 1.4. let ξ be Normal uncertain variable with expected value μ and Volatility σ^2 , Then if for any real numbers b and σ in positive real numbers $\mathbb{R}^+$, we have $|\frac{b\sigma\sqrt{3}}{\pi}| < 1$, Then

$$\mathbb{E}(e^{-bX}) = e^{-b\mu} b\sqrt{3}\sigma \csc(b\sqrt{3}\sigma) \quad (3)$$

Proof.

$$\begin{aligned} \mathbb{E}(e^{-bX}) &= \int_{-\infty}^{\infty} e^{-bx} d\left(\frac{1}{1 + e^{\frac{\pi(\mu-x)}{\sqrt{3}\sigma}}}\right) \\ &= \int_0^1 e^{-b\left(\frac{\sqrt{3}\sigma}{\pi} \ln\left(\frac{u}{1-u}\right) + \mu\right)} du \quad \left\{u = \left(1 + e^{\frac{\pi(\mu-x)}{\sqrt{3}\sigma}}\right)^{-1}\right\} \\ &= e^{-b\mu} \int_0^1 \left(\frac{u}{1-u}\right)^{-b\frac{\sqrt{3}\sigma}{\pi}} du = e^{-b\mu} \pi \left(\frac{b\sqrt{3}\sigma}{\pi}\right) \csc\left(\pi \frac{b\sqrt{3}\sigma}{\pi}\right) \\ &= e^{-b\mu} b\sqrt{3}\sigma \csc(b\sqrt{3}\sigma) \end{aligned} \quad (4)$$

□

1.2 optimization

Suppose that there are p -branch assets, $\mathbf{S} = (s_1, s_2, \dots, s_p)$ whose returns are denoted by $\mathbf{r} = (r_1, r_2, \dots, r_p)$, with mean $\mu = (\mu_1, \mu_2, \dots, \mu_p)$ and covariance matrix $\Sigma = (\sigma_{i,j})$, In addition, we suppose that an investor will invest capital $\mathbf{C}$ on the p -branch assets $\mathbf{S}$ such that person wants to allocate her or his investable wealth on the assets but obtain any of the following:

1. to maximize return subject to a given level of risk, or
2. to minimize her or his risk for a given level of expected return.

Since the above two problems are equivalent, we only look for a solution to the first problem in this paper. The target is find wight w_i of investment s_i for $1 \leq i \leq p$ such that the portfolio $W = (w_1, w_2, \dots, w_p)$ are true in one of two condition above, In other words, the return of portfolio are

$$R = \sum_{i=1}^p w_i r_i, \quad \mathbb{E}(R) = \sum_{i=1}^p w_i \mu_i \quad \text{and} \quad \text{Var}(R) = w' \Sigma w. \quad (5)$$

Definition 1.5. Exponential Utility function define

$$u(x) = 1 - e^{-bx}, \quad (6)$$

where $b \in \mathbb{R}$ are constant risk aversion.

2 Main results

In theorem (1.4) we have the condition $|\frac{b\sigma\sqrt{3}}{\pi}| < 1$. hence $b < \frac{\pi}{\sigma\sqrt{3}}$ that's mean The upper band of constant risk aversion of person symmetrical of converse volatility environment. hence, when constant risk aversion is large that's mean the person like the choice the environment with small volatility, and converse. In another side. we have $\sigma < \frac{\pi}{b\sqrt{3}}$, similarity above, if the person is small constant risk aversion can choice the environment with big volatility such that smaller than $\frac{\pi}{b\sqrt{3}}$,

if we put R (9) in utility function (6), then we have

$$u(R) = 1 - e^{-bR} \quad (7)$$

we want find the beat wight $W = (w_1, w_2, \dots, w_p)$ such that the expected of utility are maximum ($\max \mathbb{E}(u(R))$), according to the Theorem 1.4 we have,

$$\begin{aligned} \mathbb{E}(u(R)) &= 1 - \mathbb{E}(e^{-bR}) \\ &= 1 - e^{-b\mu} b\sqrt{3}\sigma \csc(b\sqrt{3}\sigma) \end{aligned} \quad (8)$$

In the equation (9) we suppose the each environment are independent of another environment's hence

$$\sigma^2 = \text{Var}(R) = w' \Sigma w = \sum_{i=1}^p w_i^2 \sigma_i^2. \quad (9)$$

hence the maximum of $\mathbb{E}(u(R))$ equal minimum of

$$f(w_1, w_2, \dots, w_n) = \frac{e^{-b\mu} b\sqrt{3}\sigma}{\sin(b\sqrt{3}\sigma)} \quad (10)$$

for find the answer we use fmincome function in MATLAB application. **Result**

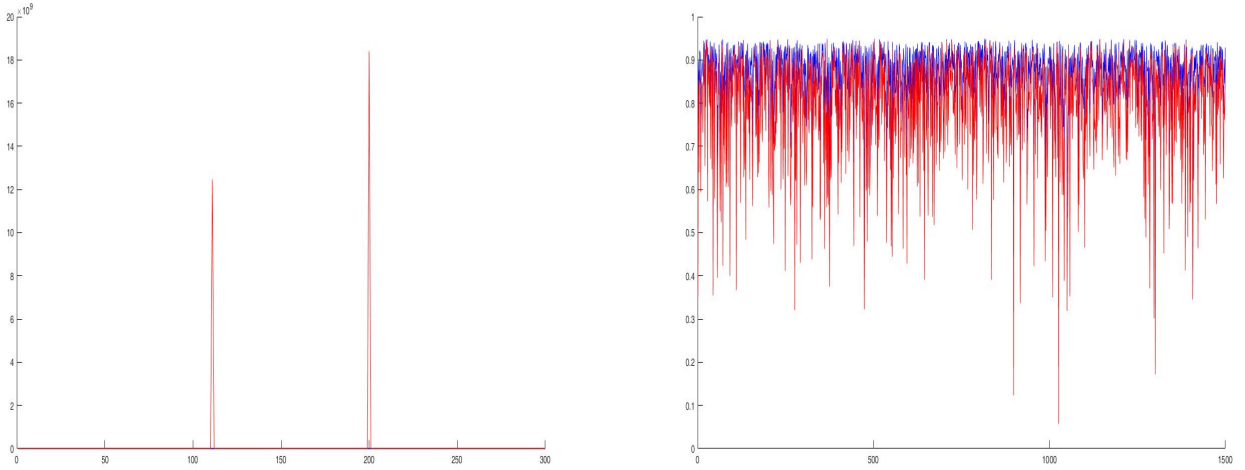


Figure 1: 1500 points of expected and volatility different with first condition of theorem (2.1) , and 2 times in 300 points the second condition in (2.1)

Theorem 2.1. suppose that $X \sim N(\mu, \sigma^2)$ is Random Normal Variable, $\xi \sim \mathbf{N}(\mu, \sigma^2)$ is Uncertain Normal Variable and $b \geq 0$ are constant risk aversion of a person-environment, also $u(x)$ is exponential utility function, then we have:

- if $b\sigma\sqrt{3} < \pi$ then

$$\mathbb{E}(u(X)) > \mathbb{E}(u(\xi)) \quad (11)$$

- if $b\sigma\sqrt{3} > \pi$ then

$$\mathbb{E}(u(X)) < \mathbb{E}(u(\xi)) \quad (12)$$

(In similar condition) The value Expected utility of Return portfolio for uncertain Normal variable is less than Expected utility with Normal variable. the sample of simulate exist in figure (??).

References

- [1] Liu, B.: *Uncertainty Theory*, 2nd edn. Springer, Berlin (2007)
- [2] Gao:1972 Gao R. (2019), *Stability in mean for uncertain differential equation with jumps*, China,.
- [3] Liu:2015 Liu, B. (2015), *Uncertainty Theory*, *Journal of the Springer Uncertainty Research*, **81**, 146–147.

e-mail: sadegh.miri@aut.ac.ir

e-mail: erfan.salavati@aut.ac.ir

The main microstructural components of the stock liquidity and intraday patterns

Nastaran Hadi Doolabi¹
Tarbiat Modares University, Tehran, Iran

Mohammad Ali Rastegar
Tarbiat Modares University, Tehran, Iran

Abstract

The aim of this study is to identify the most important liquidity measures using the intraday data of some stocks from Tehran Stock Exchange. For this purpose, the distribution features and correlation structure of a number of liquidity measures will be examined and then, using the Principal Components Analysis, these components which include the largest variances are specified and their intraday patterns will be extracted. The results show that the reduction of the number of measures to four final measures (Turnover, Liquidity Ratio², Near Market Depth and Relative Spread with mid quoted prices) is possible. Among them, the Relative Spread explanations the highest percentage of dispersion and based on the intraday patterns it is minimized in the middle of the day, so liquidity is high at these hours and favorable conditions for trading are available.

Keywords: Liquidity Measures, Market Microstructure, Intra-day Patterns, Principal Component Analysis

Mathematics Subject Classification [2018]: D47, G19, E44

1 Introduction

One of the most important parameters in deciding whether to invest in financial markets is the liquidity of different assets. Liquidity is a qualitative concept which means the ability to absorb buy and sell orders. This concept is applicable to all markets, focusing solely on liquidity in the stock market, and in particular the Tehran Stock Exchange.

In recent years, many researchers have attempted to quantify this concept and have introduced several criteria to measure it, but liquidity is a multidimensional concept that cannot be measured by a single criterion [3]. Therefore, researchers have defined four different aspects of liquidity: market depth (the effect of high volume orders on price), market width (difference between bid and ask prices), resiliency (market's ability to bounce back from temporarily incorrect prices) and Time Dimension (the speed of trades) [4]. In the current study, 27 liquidity measures have been used to quantify the liquidity of 7 selected stocks from Tehran Stock Exchange. These measures and stocks are listed in table 1 and table 2 respectively.

The data covers 77 trading days from September 22, 2016 until February 18, 2017 for 7 stocks in table 2 which are selected based on statistical reports available on the Tehran Stock Exchange website publishing the list of top 50 companies every 3 months. The data used in this study includes the intraday data from transaction prices and volumes and also the data related to limit orders available on the Limit Order Book (LOB). Over the inhomogeneous time series, a 15 second grid was imposed to get homogeneous ones with a regular spacing from 9:03:30 to 12:30:00 pm using the previous tick method. The reason for using the previous tick approach to linear interpolation is that the linear interpolation method uses future information, but the previous tick method relies solely on information up to the present.

¹ speaker

Table 1: Mathematical description of liquidity measures used in this study

Liquidity Measures	Formula	Liquidity Measures	Formula
Trading Volume	$TV_t = \sum_{i=1}^{N_t} q_i$	Relative Spread of Log Prices	$LogRelLogSp_t = \ln(\ln(p_t^A/p_t^B))$
Trading Volume per Trade	$TV_{per} = \frac{\sum_{i=1}^{N_t} q_i}{N_t}$	Effective Spread	$EffSp_t = p_t - p_t^M $
Turnover	$Turnover_t = \sum_{i=1}^{N_t} p_i \cdot q_i$	Effective Spread with Last Trade	$RelEffSp_t = p_t - p_t^M /p_t$
Number of Trades per Time Unit	N_t	Effective Spread with Mid Price	$RelEffSpMid_t = p_t - p_t^M /p_t^M$
Volume Duration	$Vdur_t = \text{Min}(Vdur: TV_{t+Vdur} \geq TV_t + V)$	Quote Slope	$Qslope_t = p_t^A - p_t^B / \ln(q_t^A) + \ln(q_t^B)$
Market Depth	$Depth_t = q_t^A + q_t^B / 2$	Log Quote Slope	$LogQslope_t = \ln(p_t^A) - \ln(p_t^B) / \ln(q_t^A) + \ln(q_t^B)$
Log Depth	$LogDepth_t = \ln(q_t^A) + \ln(q_t^B) = \ln(q_t^A \cdot q_t^B)$	Adjusted Log Quote Slope	$AdjLogQslope_t = LogQslope_t \cdot (1 + \ln(q_t^A) - \ln(q_t^B))$
Dollar Depth	$Ddepth_t = (q_t^A \cdot p_t^A + q_t^B \cdot p_t^B) / 2$	Composite Liquidity	$CL_t = RelSp/Ddepth$
Near Depth	$Ndepth_t = \frac{(\sum_{i=1}^3 q_{t,i}^A) + (\sum_{i=1}^3 q_{t,i}^B)}{2}$	Liquidity Ratio 1	$LR1_t = Turnover_t / r_t $
Near Depth Value	$NdepthV_t = \frac{(\sum_{i=1}^3 p_{t,i}^A \cdot q_{t,i}^A) + (\sum_{i=1}^3 p_{t,i}^B \cdot q_{t,i}^B)}{2}$	Liquidity Ratio 2	$LR2_t = \frac{\sum_{i=1}^{N_t} r_i }{N_t}$
Absolute Spread	$AbsSp_t = p_t^A - p_t^B$	Flow Ratio	$FR_t = N_t \cdot Turnover_t$
Log Absolute Spread	$LogAbsSp_t = \ln(p_t^A - p_t^B)$	Order Ratio	$OR_t = q_t^A - q_t^B / Turnover_t$
Relative Spread with Mid Price	$RelSp_t = p_t^A - p_t^B / p_t^M$	Order Imbalance	$\%ORI_t = \frac{TV_t^{sell} - TV_t^{buy}}{TV_t} * 100$
Relative Spread with Last Trade	$RelSpt_t = p_t^A - p_t^B / p_t$		

Table 2: List of stocks from Tehran Stock Exchange used in this study

Company Name	Persian Symbol	English Symbol	Industry
Mobarakeh Steel	فولاد	FOLD	Base Metals
I.N.C Ind.	فملی	MSMI	Base Metals
Parsian Oil&Gas	پارسان	PASN	Chemicals
Iran Khodro	خودرو	IKCO	Automotive
Metals and Min	ومعادن	MADN	Metal Ore
Ir.Inv.Petr	پترویل	IPTR	Chemicals
Saipa Inv	وسایا	SSAP	Financing

PCA Analysis. This method was developed to investigate the relationship between several variables and reduce their complexity. In this method, the variables in a multi-state correlated environment are summarized as a set of uncorrelated variables that can explain the dynamics of these variables. The obtained uncorrelated components are called principal components, each of which is derived from the linear combination of n main variable of the problem. The first principal component (PC_1) explains the highest amount of data dispersion in the entire dataset. Also, the coefficients b_{1j} are the elements of Eigenvector b_1 , which is the Eigenvector with the highest eigenvalue of the variance-covariance matrix.

$$PC_1 = b_{11} \cdot m_{11} + b_{12} \cdot m_{12} + \dots + b_{1n} \cdot m_{1n}$$

2 Main results

In this section, the results of the descriptive statistics calculations, correlation analysis, principal component analysis are presented and finally the intraday patterns are extracted.

Descriptive Statistics

In this section, mean, median, maximum, minimum, standard deviation, skewness, kurtosis, and different percentiles are calculated for each measure and per stock. General remarks about the summary statistics of the different liquidity measures across the 7 stocks are presented below:

In general, none of the measures have a symmetric distribution and are mostly skewed to the right. They are also fat-tailed in comparison to the normal distribution because of their large amounts of kurtosis.

Among the market depth-related measures, the “Log Depth” showed better distributive properties than “Market Depth” and is almost symmetric.

For measures related to the spread, the “Relative Spread of Log Prices” is better distributed than the other measures. In the case of quoted slope measures, “Log Quote Slope” with less kurtosis has a smoother distribution shape, and because of less variance, the data distribution has a higher average symmetry than the “Quoted Spread”.

Correlation Analysis

In this section, for each stock, the “Pearson Correlation Coefficient” between each 2 measures is calculated and then the average correlation of this pair is calculated over 7 stocks. According to these results the measures that are highly correlated with each other (more than 0.995) are as follows: Trading Volume and Turnover, Market Depth and Dollar Depth, Near Depth and Near Depth Value, Absolute Spread and Relative Spread with Mid Price, Relative Spread with Last Trade and Relative Spread with Mid Price, Effective Spread and Effective Spread with Last Trade, Effective Spread with Mid Price and Effective Spread with Last Trade, Log Absolute Spread and Relative Spread of Log Prices and finally Log Quote Slope and Adjusted Log Quote Slope.

Reduction of the number of Liquidity Measures

According to the results of previous sections the following measures are excluded from the study:

Trading Volume, which is highly correlated with Turnover, and the latter measure is comparable among different stocks with different prices. Number of Trades per Time Unit, which is directly used in the calculation of Volume Duration measure. Market Depth which is highly correlated with Dollar Depth, and the latter measure is comparable among different stocks with different prices. Near Depth which is highly correlated with Near Depth Value, and the latter measure is comparable among different stocks with different prices. Dollar Depth which is highly correlated with Near Depth Value, and the latter measure is more precise due to using the information of all three levels. Absolute Spread, is highly correlated with many spread measures and because relative Spreads are more comparable across stocks due to price considerations. Relative Spread with Last Trade which is perfectly correlated with Relative Spread with Mid Price, and the latter is easier to calculate because of not depending on last trade information which may not be always available. Effective Spread which is highly correlated with Effective Spread with Last Trade, and the latter measure is comparable among different stocks with different prices. Effective Spread with Last Trade which is perfectly correlated with Effective Spread with Mid Price, which is easier to calculate. Log Quote Slope which is perfectly correlated with Adjusted Log Quote Slope, and the latter is a more comprehensive measure.

Principal Component Analysis

After removing 10 measures in the last section, the principal components analysis will be performed on the remaining 17 measures. The output contains the eigenvalues, eigenvector coefficients and percentages of variance explained by each component, arranged in descending order. For all the stocks, it was observed that only 5 principal components in each stock had eigenvalues greater than one, which explain about 70 to 75 percent of the total variance. An example of the results of this analysis for stock SSAP is shown in Table 3.

Table 3: Principal Component Analysis of 17 liquidity measures for SSAP

	b_1	b_2	b_3	b_4	b_5
Eigenvalue	5.61	2.57	1.75	1.58	1.01
Var. explained (%)	33.01	15.11	10.29	9.30	5.93
Cum. Var. explained (%)	33.01	48.12	58.41	67.71	73.64

Based on the principal component analysis method performed on the 7 stocks, the following results are obtained:

One factor, explains spread-related liquidity measures. This component with the highest percentage of variance explaining across all stocks - about 33% - covers the width dimension of liquidity. From this group, the relative spread with mid-price is chosen due to the ease of calculation.

A second factor captures liquidity measures related to volume and the number of transactions. As a result, this component explains the volume and timing dimension of liquidity with a variance of about 14%. The turnover measure is considered to be representative of this group because of the ease comparability between different stocks.

The next principal component explains the measures related to the market depth, such as Log Depth. So, this factor then describes the depth dimension of liquidity explaining about 10% of variance. From this group, the near depth value measure is chosen because of the use of information on all three levels of the Limit order book.

The fourth component shows the measure related to the market resiliency dimension, such as order imbalance and the liquidity ratio2, explaining about 9% of variance. The Liquidity Ratio 2 is considered to represent this group because of its simpler interpretation.

Intraday Patterns

In this section, the Intraday behavioral pattern is extracted for the 4 selected liquidity measures in the previous section, each of which was selected as a representative of one of the liquidity dimensions.

Relative Spread with Mid Price

Various patterns including U-shaped and the inverted S-shaped patterns have been reported in the literature for this measure [1]. Here, as can be seen in the Figure 2, this measure has an inverted J-shaped pattern for most stocks, especially for IPTR and SSAP. This means that the spread is very high at the beginning of the day, but over time and with the increase in market depth, this measure falls and then rises slightly at the end of the day.

Turnover

For this measure, in most studies, the U-shaped pattern is obtained. But as can be seen in Figure 3, only IKCO partially follow this pattern. Also, the stocks FOLD and MADN are almost J-shaped. In the case of the SSAP stock, there was no clear pattern due to the sharp volatility of trading volume over the 15-second intervals.

Near Depth Value

As shown in Figure 4, various patterns for this measure have been obtained in different stocks, among which the FOLD and IPTR stocks have the S-shaped pattern. This means that at the beginning of the day, the market depth is low and at the end of the day relatively high. SSAP and IKCO have inverted U-shaped pattern.

Liquidity Ratio 2

This ratio indicates the average percentage of price changes after each transaction. As can be seen in Figure 5 the percentile graphs fluctuate strongly in all stocks, but for the mean, the graph with the lowest fluctuation is slightly above zero. No specific pattern is visible.

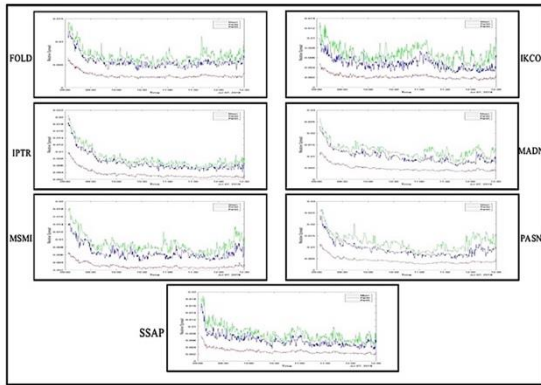


Figure 1: Intraday patterns for Relative Spread with Mid Price

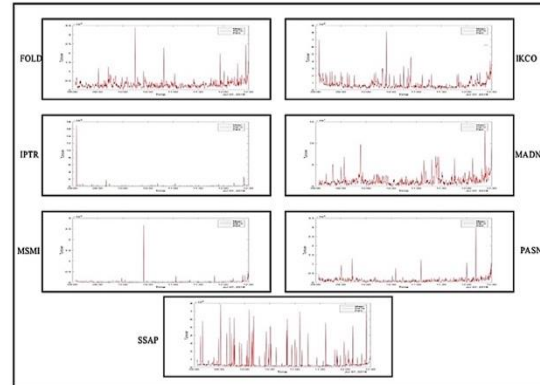


Figure 2: Intraday patterns for Turnover

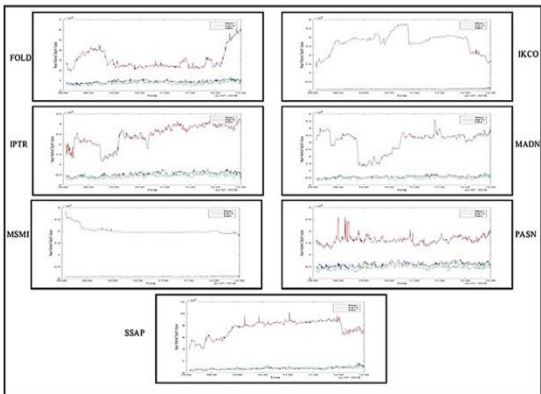


Figure 3: Intraday patterns for Near Depth Value

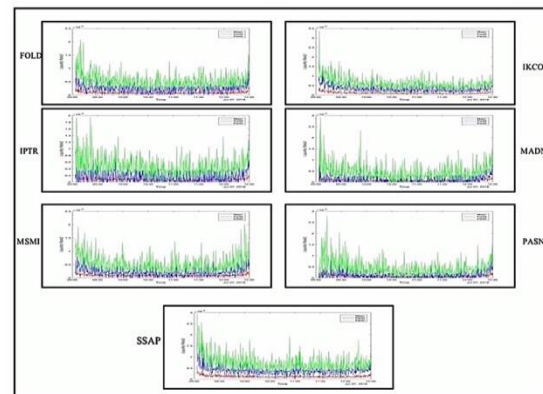


Figure 4: Intraday patterns for Liquidity Ratio 2

References

- [1] W. Cui, *An empirical investigation of price impact: An agent-based modelling approach*, Doctoral dissertation, Michael Smurfit School of Business University College Dublin, 2012.
- [2] J. Olbrys, M. Mursztyn, *Measurement of stock market liquidity supported by an algorithm inferring The initiator of a trade*. Operations Research and Decisions, 27(4): 111-127, 2017.
- [3] A. Salighehdar, Y. Liu, D. Bozdog, and I. Florescu, *Cluster analysis of liquidity measures in a stock market using high frequency data*. Journal of management science and business intelligence, 2-2, 2017
- [4] R. Wyss, *Measuring and predicting liquidity in the stock market*, Doctoral dissertation, St. Gallen University, 2004.

Email: nastaran.hadi@modares.ac.ir

Email: ma_rastegar@modares.ac.ir

Application of Factor Analysis in Assessing the Service Quality of Departments of Claim Assessment in Insurance Companies

Somayeh Mireh¹

Researcher of Insurance Research center, Tehran, Iran

Farbod Khanizadeh

Assistant Professor of Insurance Research center, Tehran, Iran

Abstract

The present paper discusses a practical approach to improve the process of claim assessment and payment procedure in Iran insurance industry. We employed factor analysis to design a proper questionnaire addressing factors affecting the dissatisfaction of insured. Then we analyzed the reasons behind insured dissatisfaction with claim assessment and payment processes. Having identified and analyzed the reasons, we combined the results with the insurer's viewpoints to identify the final factors. Finally, we provided a model suitable for insurance damage assessment and payment process.

Keywords: Factor Analysis, Service Quality, Claim Assessment, Insurance Company.

Mathematics Subject Classification [2018]: 62H25, 97M30

1 Introduction

Due to the importance of the loss assessment process, considerable research has been devoted to this area. In Nigeria, Tajudin and Adebawale (Tajudeen and Adebawale, 2013) have examined the role of managers in the loss process assessment in insurance companies. According to Marquis (2011), insurance claims management includes all company policies and insurance industry guidelines that companies use to accept or reject claims. Butler and Francis (2010) say that compensation payment is a process that indicates whether the annual premium received by insurers may be sufficient to pay the reported damages or not. From a commercial point of view, they believe that compensation payment is considered to be the biggest cost to insurers (about 80 percent of the premium).

This article focuses on the need to pay attention to customer satisfaction in the insurance industry and examines the factors affecting this phenomenon. In fact customer's satisfaction and service quality are considered as vital affairs in mostly service industry nowadays (Kuo et. al. 2009).

2 The concept of service quality

Service quality was conceptualized by the consumers as an overall assessment of service. It is believed that perceived service quality is the result of comparing the prior customer expectations for the service after actual experience with their perceptions.

The ability of a company to serve the needs of the customer and maintain its competitive advantage also affects the perception of the quality of service by the customer (Ganguli 2010). A review of studies in this area shows that

¹ speaker

the most commonly used model in measuring service quality and determining the gap between the status quo and the desired one, is the SERVQUAL model. This model was first proposed by Parasuraman et al. (Parasuraman, 1988).

The SERVQUAL model attempts to identify the major activities of the organization that affect the quality of service. Recognizing and reinforcing these factors in the insurance companies not only maintain the current insured but also attract new insured. The SERVQUAL model measures customer perceptions and expectations of service quality based on 31 components (six dimensions). The six dimensions are: Tangibles, Reliability, Responsiveness, Assurance, Empathy and Accessibility.

3 Methodology

The statistical population of the study consisted of customers who faced some loss and are referred to the automobile loss assessment and payment unit. To calculate the sample size, an initial sample of 35 was considered of which 32% were satisfied with the quality of service. With a 5% error and 95% confidence level, the required sample size can be obtained through the Cochran's formula as follows (Hekmatpo, 2012):

$$n = \frac{z_{1-\alpha/2}^2 pq}{d^2} = \frac{1.96^2 \times 0.32 \times 0.68}{0.05^2} \cong 335$$

As it is advised, more questionnaires (25% more) were distributed and finally 337 completed ones were received.

The questionnaire was designed into two parts in order to collect relevant information. The first part includes questions regarding demographic information and the second part examines the quality of service, which measures the perceptions and expectations of customers who have faced losses. For this purpose, the standard questionnaire namely; SERVQUAL was used and by means of factor analysis in LISREL and SPSS software, the scale for service quality dimensions was designed. The questions are carefully designed to ensure that they are simple and clear. The questionnaire consists of 29 questions. Furthermore 6 questions were asked at the end of the questionnaire to assess the importance of each dimension of service quality from customers' point of view.

The validity of questionnaire is verified based on content method and, using Cronbach's alpha, their reliability for expectation, perception and total questionnaire scores was 0.966, 0.921 and 0.949, respectively. Cronbach's alpha is a measure of internal consistency, that is, how closely related a set of items are as a group. In other words, the higher the α coefficient (close to 1), the more the items have shared covariance and probably measure the same underlying concept.

4 Main results

As mentioned, the purpose of this article is to measure the level of customer satisfaction using the well know instrument namely; SERVQUAL model. Basically the model is a standard questionnaire consists of different questions assessing customer satisfaction. Questions display different level of importance and some have a greater impact on people's content. Weighted average was used to analyze satisfaction weights were calculated through running Confirmatory Factor Analysis (CFA) in LISREL.

The following path diagram in LISREL shows the outcome of CFA for the questionnaire in two modes: T-values and standard factor loadings. T-values are all greater than 1.96. As a result, all questions play an important role in measuring customer satisfaction. The weighted average of codes obtained from questionnaire and factor loadings is also used to measure customer satisfaction.

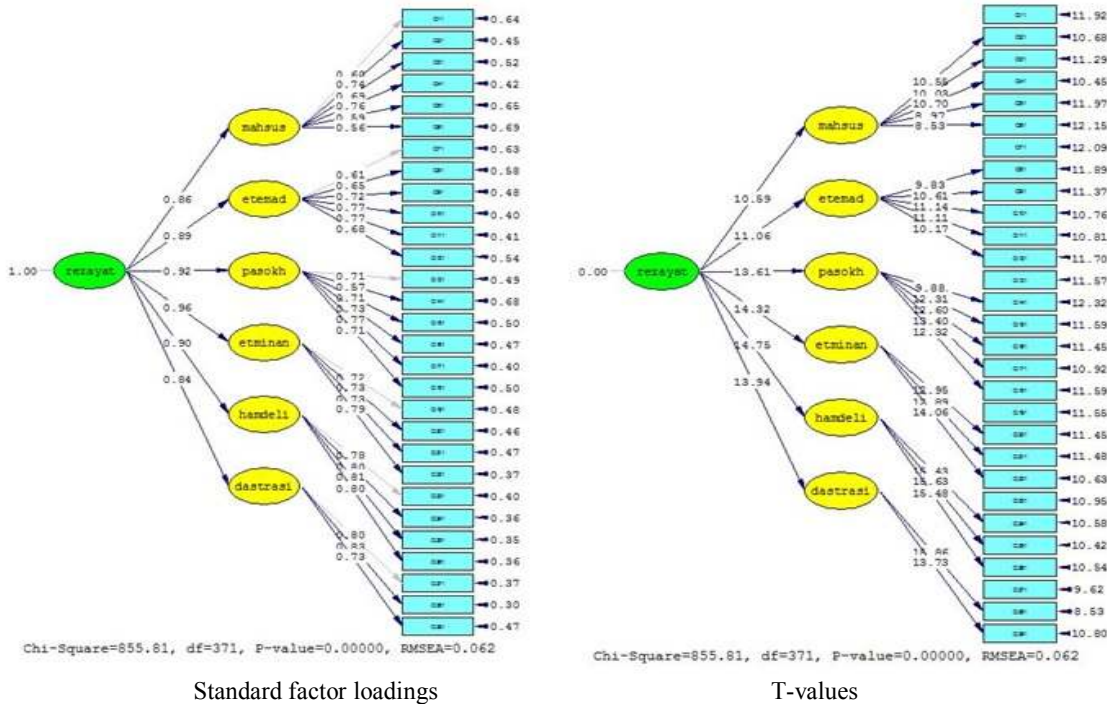
After collecting the data, we analyzed them with SPSS software using descriptive and inferential statistics such as normality test, Spearman correlation coefficients and nonparametric tests.

The descriptive statistics results table below (Table 1) indicates that insured perceive highest satisfaction in the empathy area and lowest in the tangibles area. In order to achieve higher levels of service quality, the insurance company managers should redesign their strategies about customer satisfaction with respect to service quality.

Table 1: Insured satisfaction

Dimensions of service quality	Insured satisfaction
Tangibles	0.68
Reliability	0.74
Responsiveness	0.69
Assurance	0.75
Empathy	0.83
Accessibility	0.76

Figure 1: Confirmatory factor analysis



Spearman's correlation coefficients between service quality dimensions (Table 2) also show a high correlation between reliability and tangibles of employees, which indicates that the more employees respond to the customer, the more they trust the insurance company. There is little correlation between accessibility and responsiveness of insurance company employees. Overall, there is no very high correlation between the dimensions of service quality.

Table 2: Spearman correlation coefficients of dimensions

	Tangibles	Reliability	Responsiveness	Assurance	Empathy	Accessibility
Tangibles	1	0.618	0.598	0.549	0.479	0.469
Reliability	0.618	1	0.666	0.514	0.547	0.477
Responsiveness	0.598	0.666	1	0.577	0.558	0.461
Assurance	0.549	0.514	0.577	1	0.610	0.572
Empathy	0.479	0.547	0.558	0.610	1	0.520
Accessibility	0.469	0.477	0.461	0.572	0.520	1

The results of inferential statistics (hypothesis testing) show that most customers expect that claim assessment and payment units to be available and reliable. They were also more satisfied with the tangibles of assessment units and the level of responsiveness of employees. This research also pointed out that gender, marital status, age, job, and education also had little effect on satisfaction with the assessment.

5 Conclusion

As discussed in the article, the service quality of insurance companies, especially in the area of loss assessment, is very important. The results have shown that insurers in the loss assessment unit should make more efforts to satisfy the insured.

According to the results, the following suggestions are offered to improve the quality of insurance companies' services in the claim assessment and payment unit:

- Upgrading queue management system to save customers time and money,
- Responding to customer queries on holidays and non-business hours,
- Identifying the customers' hidden needs,
- Training and encouragement of insurance company employees in dealing with customers and solving their problems.

References

- [1] S. Butler, and P. Francis. *Cutting the Cost of Insurance Claims, taking control of the process*. (2010), Booz & Co. Retrieved from www.booz.com/media/file/cutting_the_cost_of_insurance_claims.pdf
- [2] S.H. Ganguli, and S.K. Roy. *Service quality dimensions of hybrid services*. Managing Service Quality journal, 20 (2010), No.5, pp. 404-424.
- [3] D. Hekmatpo, M. Sorani, A. Farazi, Z. Fallahi, and B. Lashgarara. *A survey on the quality of medical services in teaching hospitals of Arak University of Medical Sciences with SERVQUL model in Arak*. Arak Medical University Journal; 15 (2012), No.66, pp. 1-9.
- [4] Y.W. Kuo, C.M. Wu, and W.J. Deng. *The relationships among service quality value, Customer satisfaction and post – purchase intention in mobile value – added services*. The Netherlands Computers in Human Behavior, 25 (2009), No.4, pp. 887-896.
- [5] C. Marquis. *Importance of claims management in insurance sector*.(2011), Retrieved from http://www.ehow.com/facts_7187579_importance-claims-management-insurance-sector.html
- [6] A. Parasuraman, V. Zeithaml, and L.L. Berry. *SERVQUAL: A multiple-item scale for measuring consumer perceptions of service quality*. Journal of Retailing, 63 (1988) No.1, pp.12–37.
- [7] A. Tajudeen, and A. Adebawale. *Investigating the Roles of Claims Manager in Claims Handling Process in the Nigeria Insurance Industry*, journal of Business and Finance, (2013), ISSN: pp. 2305-1825 (online), pp. 2308 - 7714 (Print).

Email: s_mireh@sbu.ac.ir

Email: khanizadeh@irc.ac.ir

Asset correlation and distance to default relationship, a case study of Iran

Mahmoud Botshekan¹

Assistant professor, University of Isfahan, Isfahan, Iran

Somayeh Mohammadi²

Ph.D. student, University of Isfahan, Isfahan, Iran

Ziba Ghazavi

MA in financial management graduated from University of Isfahan, Isfahan, Iran

Abstract

This paper examines the relationship between asset correlation and distance to default (as a proxy for probability of default) assumed in the internal ratings-based (IRB) approach of the Basel II accord on regulatory capital requirement. Using data of non-financial companies listed in Tehran stock exchange during 1388 to 1397, we show that Basel II accord (Basel, 2006) assumption on a negative relationship between asset correlation and probability of default is confirmed among TSE listed companies. Results also confirm time varying and cyclical nature of the co-variation of these two variables. Moreover, the general assumption underlying Basel II that leads to size adjustment and different equations for determining asset correlation for corporates, SMEs, and retail debtors is also consistent with the general pattern of the relationship between asset correlation and probability of default among these three categories in Iranian companies.

Keywords: Asset correlation, Distance to Default, ASRF, SME

1 Introduction

By proposing Internal Rating Based approach in Basel II accord (Basel, 2006), banks were able to determine credit risk regulatory capital charge using foundation or advanced Internal Rating Based approaches (FIRB and AIRB approaches)³ To apply these two approaches, banks need to provide four components to estimate the risk contribution of each asset to regulatory capital. These components are probability of default (PD), loss given default (LGD), exposure at default (EAD) and maturity (Basel, 2006, paragraph, 211). If banks follow foundation approach, they need to provide their own estimates of PD and can rely on supervisory general estimates for other risk components. But in the case of using advanced approach, banks also need to estimate other variables. In both cases, banks have to estimate probability of default, conditional on possible most adverse economy condition. For this purpose, the proposed framework relies on Gordy's (2003) asymptotic single risk factor (ASRF) model. In the model, there is one common systematic risk factor that affects all obligors' asset returns and correlation between an obligor's asset return and the common factor is a key parameter in determining conditional obligor's default probability and subsequently its risk contribution to regulatory capital. To be more specific,

$$1) \quad CDP = \Phi\left(\frac{\Phi^{-1}(p_i) + \sqrt{\rho}\Phi^{-1}(0.999)}{\sqrt{1-\rho}}\right)$$

¹ Corresponding author

² speaker

³ "Subject to certain minimum conditions and disclosure requirements, banks that have received supervisory approval to use the IRB approach may rely on their own internal estimates of risk components in determining the capital requirement for a given exposure" (Basel, 2006, paragraph, 211)

Where CDP is conditional default probability, p_i is unconditional default probability, $\Phi(.)$ is the standard normal cumulative density function and $\Phi^{-1}(.)$ is the inverse of that function.

So banks need to estimate unconditional default probability and also the asset correlation for each asset in order to compute the risk contribution of each asset to regulatory capital. So one of the main parameters in determining capital requirement is the asset correlation. In this regard, Basel accord also provides a fine-tuned approach for estimating asset correlation based on two premises; first, there is a negative relation between asset correlation and default probability and second, larger companies are more diversified and more sensitive to common risk factor. Based on these two assumptions that rely on some academic studies, Basel accord also differentiates among their exposures based on size⁴ and proposes three different formulas to compute asset correlation as a function of default probability. For corporate, sovereign, and bank exposures asset correlation is calculated as following:

$$2) \quad \rho(PD) = 12\% \times \frac{1 - e^{-50 \times PD}}{1 - e^{-50}} + 24\% \times \frac{1 - (1 - e^{-50 \times PD})}{1 - e^{-50}}$$

Furthermore, for corporate exposures that have less than €50 million annual sale, firm size adjustment is made in paragraph, 273. This adjustment to SME borrowers leads to correlation formula for SME exposures as:

$$3) \quad \rho^{SME}(PD) = 12\% \times \frac{1 - e^{-50 \times PD}}{1 - e^{-50}} + 24\% \times \frac{1 - (1 - e^{-50 \times PD})}{1 - e^{-50}} - 4\% \times \left(1 - \frac{(\max(S, 5) - 5)}{45}\right)$$

Where S is total annual sale.

In paragraphs 328, 329 and 330 Basel committee recommends to use new amounts for retail exposures, for example for residential mortgage exposures correlation is constant and equals to 0.15 and for qualifying revolving retail exposures it is 0.04. For other retail exposures correlation is computed as:

$$4) \quad \rho(PD) = 3\% \times \frac{1 - e^{-35 \times PD}}{1 - e^{-35}} + 16\% \times \frac{1 - (1 - e^{-35 \times PD})}{1 - e^{-35}}$$

Our motivation in this paper is to empirically examine the two premises underlying the above formulas in Iranian capital market, namely a negative correlation between asset correlation and default probability and the significant difference between asset correlations among different size classifications. To this end we follow Lee et al. (2011) and estimate asset correlation for all non-financial companies listed in Tehran Stock Exchange (TSE) assuming the market index as a single common systematic factor causing asset correlations and test the correctitude of the above mentioned premises.

To examine these relationships, following Vasallou and Xing (2004), Bharath and Shumway (2008) and Afik et al. (2016) among others, we use KMV-Merton distance to default as a proxy for default probability. The relationship between asset correlation and default probability is examined both month by month and in aggregate, for sample companies. We also divide all listed companies in TSE based on their sale to three categories and examine the mentioned relationship separately for these categories and also test whether the difference between asset correlations, as well as default probabilities, are significant among these three categories.

Although, the overall results of this study confirm the premises of Basel Accord underlying the formula mentioned before, there is not a uniform consensus in the literature. Dullmann & Scheule (2003) assuming a one factor model,

⁴ In paragraph 215 committee mentioned that, banks must differentiate among the exposures with dissimilar underlying risk types to corporate, sovereign, bank, retail and equity.” A corporate exposure is defined as a debt obligation of a corporation, partnership, or proprietorship” (Basel, 2006, paragraph, 218). Also committee detached large corporate from SME (Small and Medium Entities) in paragraph 273.

estimated constant correlation for corporate obligors of Germany. They showed that aggregate asset correlation over all of ratings increases by size, but no explicit relationship between asset correlation and PD is observed. Dietsch and Petey (2004) studied the relationship between PD and asset correlation for a large sample of SMEs in France and Germany using one factor model. They illustrate that SMEs are riskier but their asset correlations are very low. The authors also indicate that the negative relationship of PD and AC (asset correlation) in Basel formula is not confirmed among SMEs of France and Germany. The PD and AC relationship in France's SME firms is U shaped and it is positive in Germany.

2 Methodology

2.1. Estimating distance to default as a proxy of default character of firm

In this paper, following the Naive model proposed by Bharath and Shumway (2008)⁵, the market value of firm's debt and its volatility are approximated as follows:

$$\begin{aligned} \text{naive}D &= F \\ \text{naive}\sigma_D &= 0.05 + 0.25 * \sigma_E \end{aligned} \quad 5)$$

Then the asset value is equal to equity value plus debt value and asset volatility is written as:

$$\text{naive}\sigma_V = \frac{E}{E + \text{naive}D} \sigma_E + \frac{\text{naive}D}{E + \text{naive}D} \text{naive}\sigma_D = \frac{E}{E + F} \sigma_E + \frac{F}{E + F} (0.05 + 0.25 * \sigma_E) \quad 6)$$

In order to obtain the debt with monthly frequency a linear interpolation has applied to data. Volatility of equity is estimated by daily data for past one year in each single month using Newey-West (1987) estimator to control the autocorrelation in equity returns. Then distance to default is estimated in KMV model as an approximation of firms default character.

$$DD := \frac{(V_0 - F)}{\sigma_V V_0} \quad 7)$$

2.2. Asset correlation

For the asset correlation, as mentioned above, we followed Lee et al. (2011), in a way that the market factor used as a single risk factor and beta of each asset is estimated from beta of equity. In order to estimate equity beta we used CAPM model with window length varying from 36 to 60 months according to data availability. Then asset betas are estimated by the next equation:

$$\beta_V = \beta_E \left(\frac{1}{N(d_1)} \frac{E}{V} \right) \quad 8)$$

For computing $N(d_1)$ we used fix risk free rate (0.18) as the mean of asset returns, and one year annual debt as a proxy of firm debt. Finally, the asset correlation is acquired as:

$$\rho = (\beta_V \frac{\sigma_M}{\sigma_V})^2 \quad 9)$$

2.3. Firms Classification

After asset correlations are estimated, firms in TSE are classified to corporate, SME and retail sectors based on past years sale. First we estimated yearly sales of each firm in our sample in euro in every month, then we divided our sample into 3 categories due to their euro sales.

2.4. Analysis

To examine correctness of the Basel premises on negative relationship between asset correlation and default probability, two types of regressions and two types of correlation analysis have been applied. Pooled regression and Fama-Macbeth (1973) approach regressions are used. While pooled regression ignores the time effects, Fama

Macbeth regression shows time varying nature of relation between two variables. Also rank correlation and Pierson correlation are used to compute both rank and linear correlations.

For pooled regression we have following equation:

$$10) \quad AC_i = \alpha + \lambda DD_i$$

And for Fama-Macbeth regression we have:

$$11) \quad AC_{it} = \alpha_{it} + \lambda_{it} DD_{it}$$

Then the average of coefficients are reported.

3 Results

Table 1 shows that asset correlations are significantly different in dissimilar firm classes. These calculations are done by pooled data in each class. Asset correlation has its highest value in the corporate class and the lowest in retail. Besides, the maximum distance to default belongs to corporate category and minimum distance to default could be found in retail class.

Table 1. Asset correlation and distance to default comparison among different firm classifications

Panel A: Three level Asset Correlation mean comparison					
Firm class	Mean	variance	Month-firm	t-test to compare means	t-value (one tail)
Total	0.16893	0.040371			
Corporate	0.22975	0.066749	4229	Corporate-SME	15.78593
SME	0.15462	0.040047	4142	SME-Retail	10.46507
Retail	0.11479	0.032749	3026	Corporate-Retail	25.33759
Panel B: Three level Distance to Default mean comparison					
Total	2.474574	0.994078			
Corporate	2.759404	1.224536	4592	Corporate-SME	12.04961
SME	2.496411	0.904243	4312	SME-Retail	7.370356
Retail	2.343244	0.705599	3176	Corporate-Retail	18.82265

Correlation analysis is reported via table 2. Founding of rank correlation indicates that the highest correlation between asset correlation and distance to default is in the SME class, even though, the lowest rank correlation is in retail set. Pierson correlation results confirm the founding of rank correlation.

Table 2 Correlation analysis

	Rank correlation coefficient	Pierson correlation coefficient
Total	0.4130	0.4353
Corporate	0.4594	0.4806
SME	0.4756	0.5180
Retail	0.3772	0.4164

Regression analysis is the subject of table 3. Two types of regressions are reported. Results show the positive (negative) relationship between asset correlation and distance to default (probability of default) which is more strong in SME class in both Fama-Macbeth approach and pooled regressions⁶.

⁶ Reported coefficients in Fama-Macbeth regression are the average of obtained amounts

Table 3 Regression Analysis

	All companies	Corporate	SME	Retail
Panel A: Fama-Macbeth Regressions				
α	-0.072 *** (0.017) -4.277	-0.092 *** (0.032) -2.902	-0.065 *** (0.016) -3.996	-0.010 (0.015) -0.640
$\bar{\lambda}_{DD}$	0.100 *** (0.015) 6.591	0.126 *** (0.026) 4.770	0.088 *** (0.013) 6.685	0.051 *** (0.010) 4.977
$\bar{R}^2$	0.152	0.179	0.214	0.177
Panel B: Pooled Regressions				
α	-0.036 *** (0.004) -9.079	-0.051 *** (0.009) -5.655	-0.075 *** (0.007) -10.593	-0.029 *** (0.007) -4.262
λ_{DD}	0.083 *** (0.001) 55.431	0.101 *** (0.003) 33.627	0.092 *** (0.003) 34.787	0.061 *** (0.003) 22.396
$\bar{R}^2$	0.171	0.211	0.226	0.142

4. Conclusion

Using market and accounting data from 1388-1397 this study examined the relation between asset correlations and distance to default among Tehran stock exchange non-financial companies. Our findings confirm a negative relation between default probabilities and asset correlations in Iranian firms. Moreover we found that asset correlation increases by firm size. These two outcomes confirm Basel II assumption in capital requirement.

References

- 1) Afik, Zvika, Ohad Arad, and Koresh Galil. "Using Merton model for default prediction: An empirical assessment of selected alternatives." *Journal of Empirical Finance* 35 (2016): 43-67.
- 2) Basel II, Bank For International Settlements.(2006)
- 3) Bharath, Sreedhar T., and Tyler Shumway. "Forecasting default with the Merton distance to default model." *The Review of Financial Studies* 21.3 (2008): 1339-1369.
- 4) Dietsch, Michel, and Joel Petey. "Should SME exposures be treated as retail or corporate exposures? A comparative analysis of default probabilities and asset correlations in French and German SMEs." *Journal of Banking & Finance* 28.4 (2004): 773-788.
- 5) Dullmann, Klaus, and Harald Scheule. "Asset correlation of German corporate obligors: its estimation, its drivers and implications for regulatory capital." *Basel Committee's Research Task Force Workshop on Banking and Financial Stability*. 2003.
- 6) Fama, Eugene F., and James D. MacBeth. "Risk, return, and equilibrium: Empirical tests." *Journal of political economy* 81.3 (1973): 607-636
- 7) Gordy, Michael B. "A risk-factor model foundation for ratings-based bank capital rules." *Journal of financial intermediation* 12.3 (2003): 199-232.
- 8) Lee, Shih-Cheng, Chien-Ting Lin, and Chih-Kai Yang. "The asymmetric behavior and procyclical impact of asset correlations." *Journal of Banking & Finance* 35.10 (2011): 2559-2568.
- 9) Newey, Whitney K., and Kenneth D. West. "Hypothesis testing with efficient method of moments estimation." *International Economic Review* (1987): 777-787.
- 10) Vassalou, Maria, and Yuhang Xing. "Default risk in equity returns." *The journal of finance* 59.2 (2004): 831-868.

Email: m.botshekan@ase.ui.ac.ir

Email: s.mohammadi@ase.ui.ac.ir

Email: ziba.ghazavi@gmail.com

Modeling and Evaluation Iran Mutual Funds Systemic Risk: A Conditional Value at Risk Approach

Fereshteh Shahbazin¹
Islamic Azad University, Qazvin, Iran

Hassan Qhalibaf-Asl
Alzahra University, Tehran, Iran

Mohsen Saeighali
Islamic Azad University, Qazvin, Iran

Moslem Peymani Foroushani
Allameh-Tabatabaei University, Tehran, Iran

Abstract

Mutual funds are the most important investment companies in the financial markets that invest the collected money in various markets. Although mutual funds do not have interconnection together, they are correlated through their portfolios and their investment strategies, which may cause systemic risk. In this article, we try to model and evaluate Iran mutual funds systemic risk by Conditional Value at Risk (CoVaR) approach. Results indicate share funds are riskier than other funds, say fixed-income and mixed funds. ANOVA shows a significant correlation between the funds' systemic risk and funds' portfolio beta coefficient.

Keywords: Systemic Risk, Mutual Funds, CoVaR, Quantile Regression.

Mathematics Subject Classification [2018]: 13D45, 39B42

1 Introduction

According to the most presented definitions, systemic risk refers to the possibility of financial system failure due to a crisis or big shock in a segment or some parts of the financial market. When a company defaults and is unable to achieve its goal a shock will occur in the system and other companies may default as well. This paper's aim is to research and identify systematically important mutual funds, which have a significant proportion of systemic risk, particularly. Regulators and market authorities reduce the effect of systemic risk by enacting roles and macro-prudential policies like capital requirements and liquidity restrictions.

Mutual funds are the most important investment companies in the financial markets that invest the collected money in various markets. There are three types of securities funds in Iran capital market, including share, fixed-income, and mixed funds. Although mutual funds do not have interconnection together, they are correlated through their portfolios and their investment strategies, which may cause systemic risk.

¹ speaker

In recent years, the importance of funds' role in financial markets is obvious for all market participants and regulators. A shock can impact on target markets and other parts will be influenced by this shock. In this article, we try to model and evaluate Iran mutual funds systemic risk by Conditional Value at Risk (CovaR) approach.

We use Conditional Value at Risk (CoVaR) introduced by Adrian and Brunnermeier in 2016, as a systemic risk measure, which has become popular after the 2007-2008 global crisis. VaR concentrates on the individual company and shows the maximum amount of loss in a company with the $q\%$ of confidence level, which mathematically is:

$$Prob(X^i \leq VaR_q^i) = q\%$$

Where X^i is the fund's total return:

$$X_t^i = \frac{NAV_{i,t} - NAV_{i,(t-1)} + I_t}{V_{i,(t-1)}},$$

$NAV_{i,t}$ and $NAV_{i,(t-1)}$ are net asset value of the fund i in time t and $t - 1$, respectively and I_t is interest income of the fund i in time t .

CoVaR is defined as the system's Value of Risk so that fund i is in risk:

$$Prob(X^{system} | C(X^i) \leq CoVaR_q^{system|C(X^i)}) = q\%$$

Where X^{system} is the value-weighted of funds total returns.

Quantile regression

Quantile regression is a method to calculate VaR and CoVaR. Using this approach, we can estimate the maximum loss of system in $q\%$ percentile.

$$\hat{X}_q^{system|X^i} = \hat{\alpha}_q^i + \hat{\beta}_q^i X^i$$

Where $\hat{X}_q^{system|X^i}$ is a $q\%$ -quantile estimation of the system on observed returns of the fund i . By definition, the VaR is calculated by:

$$VaR_q^{system|X^i} = \hat{X}_q^{system|X^i}$$

In practice, the system loss on fund i estimated by quantile regression is a VaR on X^i condition because $VaR_q^{system|X^i}$ is a conditional quantile. Using $X^i = VaR_q^i$, we can estimate $CoVaR_q^i$.

$$CoVaR_q^i = VaR_q^{system|X^i=VaR_q^i} = \hat{\alpha}_q^i + \hat{\beta}_q^i VaR_q^i$$

sois: $\Delta CoVaR_q^i$

$$\Delta CoVaR_q^i = CoVaR_q^i - CoVaR_{50}^i = \hat{\beta}_q^i (VaR_q^i - VaR_{50}^i)$$

$CoVaR_q^i$ indicates the system's loss when the fund i is in $q\%$ loss and $\Delta CoVaR_q^i$ shows the deterioration of the funds' system while moves from a normal state to the worst scenario.

2 Main results

The data used in this article contains total return and net asset value of 152 exchange-traded and mutual funds (shares, fixed-income and mixed) in the period between March 2016 to the end of September 2019, which has been active for 75% of the period in a weekly basis.

The result shows that Bourseiran and Goharnafis were systematically important funds at the end of Sep. 2019. Table (1) indicates the five top funds that were systemically important.

Table 1: Five top funds that were systemically important

Rank	Name	Type
1	Bourseiran	Share
2	Goharnafis	Mixed
3	Ofogh Mellat	Share
4	Bank Dey	Share
5	Tajrobe Iranian	Mixed

In the next step, in order to analyze the results more precisely, the funds are classified into 5 classes. This means that the funds in the first category are more important than the other categories in terms of systemic risk and the funds in the fifth class are the least important. This classification is performed weekly and includes variables such as return, net asset value, VAR, beta, CoVaR, and Delta CoVaR.

Since the categorization is based on the funds ranks that are based on Delta CoVaR, it is evident that the first-class funds have higher CoVaR and Delta CoVaR. Figure (1) shows the diagrams of these variables in each of the five categories.

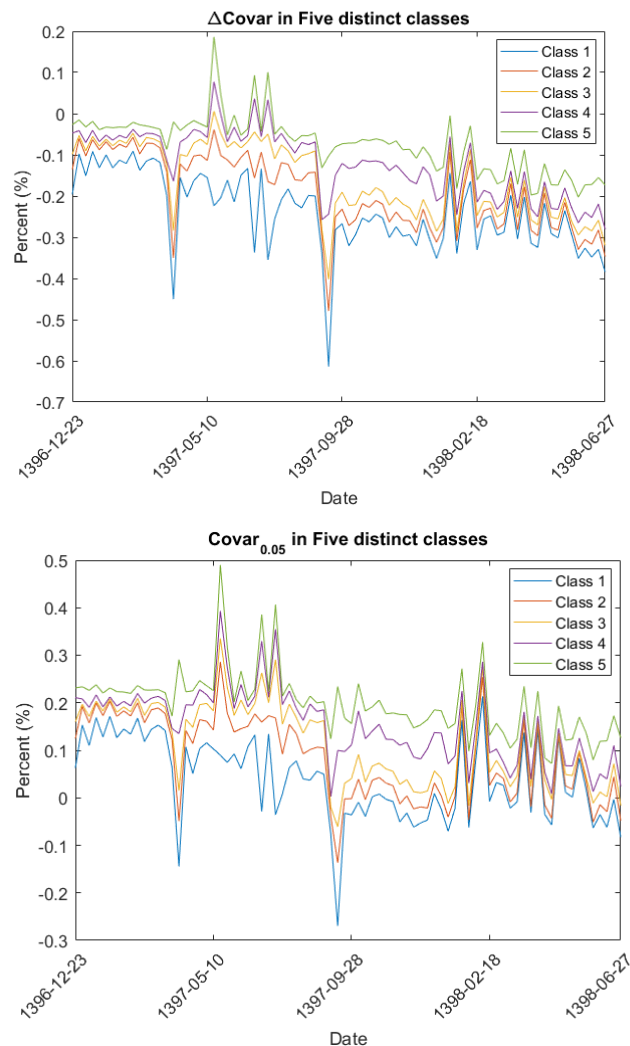


Figure 1: CoVaR and Delta CoVaR in five distinct classes

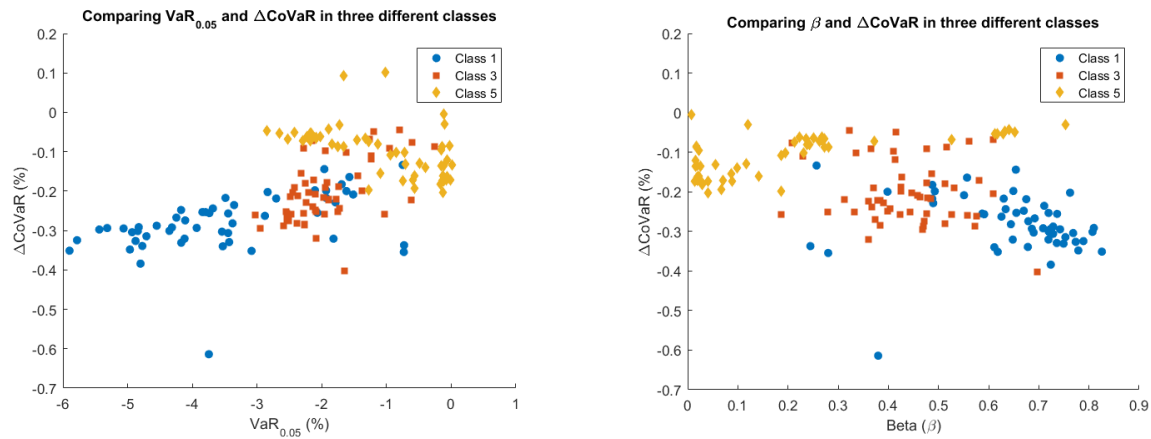


Figure 2: scatter plot beta coefficients and the risk value of the first, third and fifth classes relative to the Delta CoVaR

Given the above result, it is expected that the behavior of the parameters in each class will depend on the funds' behavior in that class. For example, the beta coefficient and the Value at Risk of funds that are in the higher classes are higher than those of the lower classes. Figure (2) shows the scatter plot beta coefficients and the risk value of the first, third and fifth classes relative to the Delta CoVaR.

ANOVA test approves the previous result and shows there are at least two classes that are different beta on average. Table (2) summarizes the result of ANOVA.

Table 2: ANOVA results

Changes	SS	df	MS	F	Prob>F
Cullum	6.812827	4	1.703207	38.10454	4.40E-25
Error	11.62155	260	0.044698		
SUM	18.43438	264			

References

- [1] Acharya, V.V., et al., Measuring systemic risk. 2010.
- [2] Acharya, V., R. Engle, and M. Richardson, Capital shortfall: A new approach to ranking and regulating systemic risks. *The American Economic Review*, 2012. 102(3): p. 59-64.
- [3] Adrian, T. and M.K. Brunnermeier, CoVaR. 2016, National Bureau of Economic Research.
- [4] Girardi, G. and A. Tolga Ergün, Systemic risk measurement: Multivariate GARCH estimation of CoVaR. *Journal of Banking & Finance*, 2013. 37(8): p. 3169-3180.

Email: f_shahbazin@yahoo.com

Bilateral Counterparty Risk Valuation to Credit Default Swaps

Sahar Salimi¹

Institute for Advanced Studies in Basic Sciences, Zanjan, Iran

Bahar Akhtari

Research fellow of workgroup financial mathematics
at Ludwig Maximilians University of Munich, Germany

Abstract

One of the most important risks, which banks, financial institutions and corporations face with, is counterparty credit risk. The counterparty credit risk is a combination of market and credit risks. One of the key tools to measure the risk of the counterparty credit is credit valuation adjustment. In this paper, under the assumption of non-arbitrage, we deal with bilateral counterparty risk valuation with stochastic dynamical models and application to credit default swaps. To do this, we utilize intensity models for default times. Then, we calculate the credit valuation adjustment using trivariate copula functions regarding the correlation between the default times and credit spread volatilities.

Keywords: Bilateral Counterparty Risk, Credit Valuation Adjustment, Credit Default Swap.

Introduction

The credit default swap (CDS) agreement transfers the default risk to the insurer. A credit default swap is a contract between the protection buyer and the protection seller that is closed on the reference credit, where the protection buyer pays periodic payments to the protection seller as a premium. In contrast, if the reference credit fails, the seller pays the fixed payment protection to the protection buyer. This type of swap contract is one of the most applicable credit derivatives, however, we can not claim that a CDS is a contract insurance but we can call the CDS by default protection [3]. An important tool for measuring counterparty credit risk is credit valuation adjustment (CVA). From the investor's point of view, unilateral credit valuation adjustment is the difference between the value of the portfolio when there is no probability of default for the counterparty with that there is a probability of default for the counterparty. The same is true for bilateral credit valuation adjustment, the difference is that a probability of default on both sides of the contract.

1 Arbitrage-Free Valuation of Bilateral Counterparty Risk

We denote by " A_1 " and " A_2 " the investor and counterparty involved in the swap financial contract where the portfolio exchanged by the two parties is default sensitive. Also we refer to the reference credit of the portfolio as " A_3 " and denote by τ_{A_1} , τ_{A_3} and τ_{A_2} respectively the default times of the investor, the reference credit and counterparty. We place ourselves in a probability space $(\Omega, \mathcal{G}, \mathcal{G}_t, \mathbb{Q})$ where filtration $\mathcal{G}_t$ models the flow of information of the whole market, including credit and $\mathbb{Q}$ is the risk neutral measure [1].

¹speaker

1.1 CDS Payoff

In a CDS contract, if A_3 , the reference credit, defaults at time $\tau = \tau_{A_3}$ with $T_a < \tau < T_b$, then A_2 , protection seller, pays protection leg to A_1 that is protection buyer. The protection leg is a certain deterministic or stochastic amount LGD_{A_3} (Loss Given Default of the reference credit " A_3 ") of the notional. In return, " A_1 " pays " A_2 " premium leg times $T_{a+1}, \dots, T_b$. The premium leg, which is a rate or a spread, is stopped as A_3 defaults. Note that the spread S is the rate of which the present value of the contract cashflows from the protection and the premium leg are equal. So protection buyer receive following cashflow at time T_j for $T_a < T_j < T_b$:

$$\begin{aligned}
CDS_{a,b}(T_j, S, LGD_{A_3}) &= Premium_{a,b}(T_j, S, LGD_{A_3}) - Protection_{a,b}(T_j, S, LGD_{A_3}) \\
&= 1_{\tau > T_j} \overline{CDS}_{a,b}(T_j, S, LGD_{A_3}) \\
&= 1_{\tau > T_j} \left\{ S \left[- \int_{\max\{T_a, T_j\}}^{T_b} D(T_j, t) (t - T_{\gamma(t)-1}) d\mathbb{Q}(\tau > t | \mathcal{G}_{T_j}) \right. \right. \\
&\quad + \sum_{i=\max\{a,j\}+1}^b D(T_j, T_i) \alpha_i \mathbb{Q}(\tau > T_i | \mathcal{G}_{T_j}) \left. \right] \\
&\quad \left. + LGD_{A_3} \left[\int_{\max\{T_a, T_j\}}^{T_b} D(T_j, t) d\mathbb{Q}(\tau > t | \mathcal{G}_{T_j}) \right] \right\},
\end{aligned} \tag{1}$$

where $t \in [T_{\gamma(t)-1}, T_{\gamma(t)})$, $\gamma(t)$ is the first premium payment date after t ($\gamma(t)$ is the first payment period T_j following time t), and α_i is the year fraction between T_{i-1} and T_i . Here, we have denoted by $D(t, T)$ the price of a zero coupon bond at t with maturity T .

Theorem 1.1. Bilateral Risk Credit Valuation Adjustment for Receiver CDS

The BR-CVA at time t for a receiver CDS contract (protection buyer) running from time T_a to time T_b with premium S is given by

$$\begin{aligned}
BR-CVA-CDS_{a,b}(t, S, LGD_{1,2,3}) &= \\
&LGD_{A_1} \cdot \mathbb{E}_t \{ 1_{\tau_{A_2} \leq \min\{\tau_{A_1}, T\}} \cdot D(t, \tau_{A_1}) \cdot [-1_{\tau > \tau_{A_1}} \overline{CDS}_{a,b}(\tau_{A_1}, S, LGD_{A_3})]^+ \} \\
&- LGD_{A_2} \cdot \mathbb{E}_t \{ 1_{\tau_{A_1} \leq \min\{\tau_{A_2}, T\}} \cdot D(t, \tau_{A_2}) \cdot [1_{\tau > \tau_{A_2}} \overline{CDS}_{a,b}(\tau_{A_2}, S, LGD_{A_3})]^+ \},
\end{aligned} \tag{2}$$

where $LGD_{1,2,3} = (LGD_{A_1}, LGD_{A_2}, LGD_{A_3})$.

1.2 Default Correlation

The default intensities of the three names A_i are denoted by $\lambda_i(t)$ for $i = 1, 2, 3$ where $\lambda_i(t)$ are independent CIR processes with cumulative intensities $\Lambda_i(t) = \int_0^t \lambda_i(s) ds$. Thus default stopping times are defined by $\tau_i = \Lambda_i^{-1}(\xi_i)$ where ξ_i are standard (unit-mean) exponential random variables. Define unit uniform random variables U_i with $U_i = 1 - \exp(-\xi_i)$. In the case of bilateral CVA, we impose a dependence structure on τ_{A_1} , τ_{A_3} and τ_{A_2} via a trivariate Gaussian copula C_Σ on U_1 , U_2 and U_3 ,

$$C_\Sigma(u_1, u_1, u_3) = \mathbb{Q}(U_1 < u_1, U_2 < u_2, U_3 < u_3), \tag{3}$$

where $\Sigma = [r_{ij}]_{i,j=1,2,3}$ is a 3-dimensional correlation matrix that parametrizes the trivariate Gaussian distribution. Consider the following stochastic intensity model for the three names A_1, A_2, A_3

$$\lambda_j(t) = y_j(t) + \psi_j(t; \beta_j), \quad t \geq 0, \quad j = 1, 2, 3, \tag{4}$$

where ψ is a deterministic function, depending on the parameter vector $\beta_j = (\kappa_j, \mu_j, v_j, y_j(0))$ with κ_j , μ_j , v_j and $y_j(0)$ positive deterministic constants. We assume that each y_j to be a Cox Ingersoll Ross (CIR) process given by

$$dy_j(t) = \kappa_j(\mu_j - y_j(t))dt + v_j \sqrt{y_j(t)} dZ_j(t), \quad j = 1, 2, 3, \tag{5}$$

with assumption that $2\kappa_j\mu_j > v_j^2$. We assume that Z_j to be a standard Brownian motion processes under the risk neutral measure. The following integrated quantities will be extensively used in the remainder of the paper

$$Y_j(t) = \int_0^t y_j(s)ds, \quad \Psi_j(t) = \int_0^t \psi_j(s; \beta_j)ds. \quad (6)$$

1.3 Calculation of Survival Probability

We recall that the survival probabilities associated with a CIR intensity process are given by

$$\mathbb{Q}(\tau_{A_i} > t) = \mathbb{E}[\exp(-Y_i(t))] = P^{CIR}(0, t, \beta_i), \quad i = 1, 2, 3, \quad (7)$$

where $P^{CIR}(0, t, \beta_i)$ is the price at time 0 of a zero coupon bond maturing at time t under a stochastic interest rate dynamics given by the CIR process, with $\beta_i = (y_i(0), \kappa_i, \mu_i, v_i)$, being the vector of CIR parameters, $i = 1, 2, 3$. The market survival probability from the definition of the integrated process $\Psi_i(t; \beta_i)$ are given by

$$\Psi_i(t; \beta_i) = \log \left(\frac{\mathbb{E}[\exp(-Y_i(t))]}{\mathbb{Q}(\tau_i > t)_{\text{market}}} \right) = \log \left(\frac{P^{CIR}(0, t, \beta_i)}{\mathbb{Q}(\tau_{A_i} > t)_{\text{market}}} \right), \quad i = 1, 2, 3. \quad (8)$$

We set $\bar{U}_{i,j} = 1 - \exp(-\Lambda_i(\tau_{A_j}))$ and denote by $F_{\Lambda_i}(t)$ the cumulative distribution function of the cumulative (shifted) intensity of the CIR process associated to name i . The credit reference survival probability at the default time of the counterparty is given by

$$\begin{aligned} 1_{C \cup D} 1_{\tau_{A_3} > \tau_{A_2}} \mathbb{Q}(\tau_{A_3} > t | \mathcal{G}_{\tau_{A_2}}) = \\ 1_{\tau_{A_2} \leq T} 1_{\tau_{A_2} \leq \tau_{A_1}} (1_{\bar{E}} + 1_{\tau_{A_2} < t} 1_{\tau_{A_3} \geq \tau_{A_2}} \int_{\bar{U}_{3,2}}^1 F_{\Lambda_3(t) - \Lambda_3(\tau_{A_2})}(-\log(1 - u_3) - \Lambda_3(\tau_{A_2})) dC_{A_3|A_1, A_2}(u_3; U_2)), \end{aligned} \quad (9)$$

and the credit reference survival probability at the default time of the investor is given by

$$\begin{aligned} 1_{E \cup B} 1_{\tau_{A_3} > \tau_{A_1}} \mathbb{Q}(\tau_{A_3} > t | \mathcal{G}_{\tau_{A_1}}) = \\ 1_{\tau_{A_1} \leq T} 1_{\tau_{A_1} \leq \tau_{A_2}} (1_{\bar{B}} + 1_{\tau_{A_1} < t} 1_{\tau_{A_3} \geq \tau_{A_1}} \int_{\bar{U}_{3,1}}^1 F_{\Lambda_3(t) - \Lambda_3(\tau_{A_1})}(-\log(1 - u_3) - \Lambda_3(\tau_{A_1})) dC_{A_3|A_2, A_1}(u_3; U_1)), \end{aligned} \quad (10)$$

where $E = \{\tau_{A_1} \leq \tau_{A_3} \leq T\}$, $B = \{\tau_{A_1} \leq T \leq \tau_{A_3}\}$, $C = \{\tau_{A_3} \leq \tau_{A_1} \leq T\}$, $D = \{\tau_{A_3} \leq T \leq \tau_{A_1}\}$, $\bar{E} = \{t \leq \tau_{A_2} \leq \tau_{A_3}\}$, $\bar{B} = \{t \leq \tau_{A_1} \leq \tau_{A_3}\}$, $C_{A_3|A_1, A_2}$ and $C_{A_3|A_2, A_1}$ are associated Gaussian conditional copula functions.

2 Example

We consider an investor (A_1) trading a five-years CDS contract on a reference credit (A_3) with a counterparty (A_2). Both the investor and the counterparty are subject to default risk. We experiment on different levels of credit risk and credit risk volatility of A_1, A_2, A_3 which are specified by the parameters of the CIR processes in Table 1. We assume that the spreads in Table 2 are the spreads quoted in the markets for A_1, A_2, A_3 under consideration.

Credit Risk Levels	y_0	κ	μ	Credit Risk Volatilities	ν
Low	0.00001	0.9	0.001	Low	0.01
Middle	0.01	0.8	0.02	Low	0.01
High	0.03	0.5	0.05	High	0.5

Table 1: The credit risk levels and credit risk volatilities parameterizing the CIR processes

Maturity	Low Risk	Middle Risk	High Risk
1 y	0	92	234
2 y	0	112	248
3 y	1	120	251
4 y	1	124	253
5 y	1	126	254

Table 2: Spreads in basis points generated using the parameters of the CIR processes.

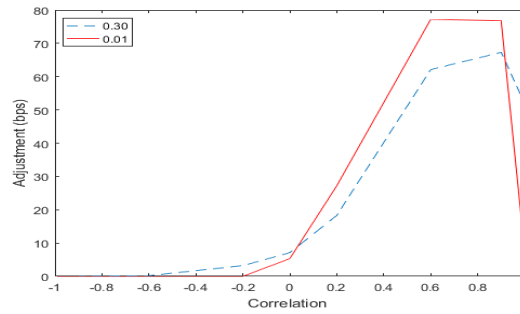


Figure 1: Bilateral Risk Counterparty Value Adjustment Payer Investor

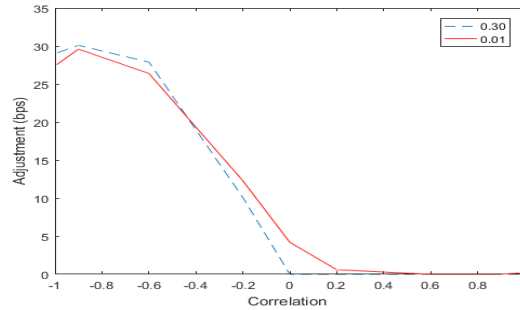


Figure 2: Bilateral Risk Counterparty Value Adjustment Receiver Investor

References

- [1] Brigo, D. and Capponi, A. Bilateral counterparty risk valuation with stochastic dynamical models and application to credit default swaps. arXiv:0812.3705, 2008.
- [2] Brigo, D. and Chourdakis, K. Counterparty risk for credit default swaps: Impact of spread volatility and default correlation. International Journal of Theoretical and Applied Finance, 12(7):1007-1026, 2009.
Brigo, D., and Masetti, M. Risk neutral pricing of counterparty risk. 2005.
- [3] Schmaltz, Ch. and Thivaos, P. Are credit default swaps credit default insurances?. Journal of Applied Business Research (JABR), 30.6 : 1819-1830, 2014.

Email: sahar.salimi92@gmail.com

Email: akhtari@math.lmu.de

Pricing of index-based catastrophe bonds Based on Fourier cosine expansions

Sahar Yaghoubi¹

Azarbaijan Shahid Madani University, Tabriz, Iran

Mojtaba Ranjbar

Azarbaijan Shahid Madani University, Tabriz, Iran

Ali Foroush Bastani

Institute for Advanced Studies in Basic Sciences, Zanjan, Iran

Abstract

We propose an efficient pricing method for catastrophe bonds based on Fourier cosine expansions. The bond is based on the PCS index which posts quarterly estimates of industry-wide hurricane losses. The aggregate PCS index is analogous to losses claimed under traditional reinsurance in that it is used to specify a reinsurance layer.

Keywords: Fourier cosine expansions, catastrophe bond, PCS index.

1 Introduction

Earthquakes, hurricanes, tornadoes and hailstorms consist of the four most costly types of insured catastrophic perils in the United States. Of these, earthquakes and hurricanes pose the greatest catastrophic risk generating on average \$9.7 billion in claims annually from 1989 through 2001. Focusing specifically on hurricanes; Hurricanes Katrina, Wilma, Rita, Ophelia and Dennis caused \$52.7 billion in insured losses in 2005 amounting to nearly 93 percent of all losses from catastrophic perils that year. In addition, Hurricanes Charley, Ivan, Frances and Jeanne produced \$23 billion in insured losses in 2004. In contrast, Hurricane Andrew alone caused \$15.5 billion in insured damages in 1992. Insurance companies often require reinsurance to limit their liabilities given the large capital requirements needed to cover these damages. Unfortunately, capital in the reinsurance industry is also limited relative to the magnitude of these damages, creating large fluctuations in the price and availability of reinsurance during years when catastrophic losses are excessive. In response, reinsurance companies have recently turned to the capital markets by issuing risk-linked securities in the form of CAT bonds ² to provide the collateral necessary for reinsurance. If a specified

¹speaker

²Catastrophe Bonds

catastrophe does not occur (or if aggregate damages are less than a trigger level) before the maturity date of the bond, the investors get the full face value of the bond plus very generous coupon payments. If the specified catastrophe does occur (or if aggregate damages exceed the trigger level) before the maturity date, the bond defaults resulting in either a partial or no payment to investors. Fortunately, the capital markets are extremely large (approximately \$31 trillion) relative to the scale of the property damage and can readily absorb this risk.

Pricing models for catastrophic risk-linked securities have followed two main financial methodologies: the theory of equilibrium pricing and the no-arbitrage valuation framework. Aase (1999), Cox and Pedersen (2000) and Cox et al. (2000) used the theory of equilibrium pricing while Sonderman (1991), Cox and Schwebach (1992), Cummins and Geman (1995), Geman and Yor (1997), Louberge et al. (1999), Lee and Yu (2002), Dassios and Jang (2003), Romaniuk (2003), Burnecki and Kukla (2003), Cox et al. (2004), and Hardle and Lopez Cabrera (2007) used the no-arbitrage framework. The primary inspiration for this research follows from Cummins and Geman (1995) who focused on pricing catastrophe insurance futures and call spreads which were traded on the Chicago Board of Trade between 1992 and 1999.

Frequency, magnitude and reinsurance layers

The frequency and magnitude of catastrophic events are arguably the two most important parameters when engineering a CAT bond contract. The focus of this work is restricted to hurricane damages along the Gulf and Atlantic coasts of the US. The frequency of landfalling hurricanes (and other catastrophic events) is generally modeled using a discrete probability density function (PDF). Typically, this PDF is characterized using a Poisson distribution where λ denotes the event frequency. Cummins and Geman (1995) state that the actuarial data indicate the event frequency of landfalling hurricanes along the Gulf and Atlantic coasts of the US is $\lambda^{\mathbb{P}} = 0.5$ per year. The superscript $\mathbb{P}$ denotes the probability measure. This value is adopted in the base-scenario CAT bond model presented here. The potential damage caused by a landfalling hurricane is estimated using a continuous PDF such as a lognormal or Pareto distribution. This means that there is uncertainty in predicting the damage due to variability in wind speeds, response of insured property to the wind load, and whether the hurricane strikes a developed area containing insured properties or not. Pielke and Landsea (1998) used a lognormal distribution to fit historical insured and uninsured hurricane damages after normalizing these data to 1995 values by adjusting for inflation, wealth and population. Accounting for the latter two factors is particularly important as development along the Gulf and Atlantic coasts has rapidly increased in recent years. The resulting annual hurricane distribution shows no trend, and Pielke (2005) concludes that the escalating cost of hurricane damages is due only to increased societal vulnerability. The base-scenario CAT bond model presented in this work utilizes, but is not restricted to, a lognormal PDF. The formula for this PDF is:

$$g(\eta) = \frac{e^{-(1/2)((\ln \eta - \mu)/\gamma)^2}}{\gamma \eta \sqrt{2\pi}}, \quad (1)$$

where for the base-scenario, $\gamma^{\mathbb{P}} = 1.11$ and $\mu^{\mathbb{P}} = -2.54$. This creates a distribution with a heavy tail, which is typical of damages caused by catastrophic events. The uncertainty in damages is converted to risk by multiplying η by the catastrophic claim amount C . For the base-scenario, $C^{\mathbb{P}} = \$40$ billion, with the scaled PDF for risk given as:

$$g^{\mathbb{P}}(\eta C^{\mathbb{P}}) = \frac{1}{C^{\mathbb{P}}} g^{\mathbb{P}}(\eta). \quad (2)$$

Therefore, the mean damage is \$3.15 billion with a 1% probability that damages will exceed \$40 billion, and a 0.4% probability that damages will exceed \$60 billion. Collateral provided through the issue of CAT bonds is typically targeted to provide reinsurance capacity against aggregate damages from catastrophic events between 1-in-100 (1%per annum) and 1-in-250 (0.4%per annum) year recurrence intervals. This layer is often not covered by traditional reinsurance because at this high severity low frequency level, buyers of reinsurance become concerned about the credit risk of the reinsurer. In addition, reinsurance premiums at this level are often uneconomical. The above definition denotes that information regarding both frequency and magnitude are needed to specify a reinsurance layer.

Formulation of the CAT bond $S - r$ component: According to [5], the formulation of the CAT bond $S - r$ component of the proposed CAT bond model is based on two stochastic processes: the PCS ³ index S and the three-month LIBOR interest rate r . To begin the derivation, the three-month LIBOR interest rate is assumed to be constant. Therefore, the only stochastic variable is the PCS index S which is assumed to follow a geometric Brownian motion (GBM) process with drift and jump diffusion (JD) as:

$$\begin{aligned} dS|_{GBM} &= \alpha S dt + \sigma_s S dW^{\mathbb{Q}}, \\ dS|_{JD} &= \eta C^{\mathbb{Q}} dq^{\mathbb{Q}}, \end{aligned} \quad (3)$$

where α is the rate of damage appreciation as measured by the PCS index over time, and includes the effects of inflation, population, wealth, and CAT coverage, σ_s represents volatility in the PCS index due to small catastrophes and randomness in reporting claims to PCS between quarters, $\eta C^{\mathbb{Q}}$ denotes the jump in the PCS index due to a large catastrophe, and $dq^{\mathbb{Q}}$ is an increment of a Poisson process where in the interval $[t, t + dt]$:

$$dq^{\mathbb{Q}} = \begin{cases} 1 & \text{with probability } \lambda^{\mathbb{Q}} dt, \\ 0 & \text{with probability } 1 - \lambda^{\mathbb{Q}} dt. \end{cases} \quad (4)$$

Now, the total change in the PCS index variable S can be represented as:

$$dS|_{total} = dS|_{GBM} + dS|_{JD}.$$

payoff structure: The proposed index-based CAT bond is a function of the augmented-state variable A which represents the aggregate PCS index. A denotes the running sum of the stochastic variable S in an analogous manner to a discrete Asian option.

$$A^n = \sum_{n=1}^M S^n,$$

where A^n is the discrete running sum obtained by sampling S at n observation times $t_1, t_2, \dots, t_M$ so that $S^n = S(t_n)$. Following the earlier notation, the lower and upper reinsurance layers are given as L_l and L_u , respectively. Now, the payoff condition is:

$$\begin{cases} \beta = 1 & A < L_l, \\ \beta = \frac{L_u - A}{L_u - L_l} & L_l < A < L_u, \\ \beta = 0 & A > L_u, \end{cases} \quad (5)$$

where

$$B(A, T) = \beta \times \text{Face Value}, \quad (6)$$

let Face Value = \$1. (more information can be found in [6]).

³Property Claim Services

2 Main results

Pricing of zero-coupon CAT bond:

Fourier integrals and cosine series: The point of departure for pricing of zero-coupon CAT bonds with numerical integration techniques is the risk-neutral valuation formula:

$$\begin{aligned} B(x, t_0) &= e^{-r\Delta t} \mathbb{E}^Q[B(A, T)|x] \\ &= e^{-r\Delta t} \int_{\mathbf{R}} B(A, T) f(A|x) dA, \end{aligned} \quad (7)$$

where $f(A|x)$ is the probability density of A given x , and r is the risk-neutral interest rate, and $\mathbb{E}^Q$ is the expectation operator under risk-neutral measure Q .

Inverse Fourier integral via cosine expansion. Since the density rapidly decays to zero as $A \rightarrow \infty$ in (7), we truncate the infinite integration range without losing significant accuracy to $[a, b] \subset R$ (more information can be found in [2]), we obtain approximation B_1 :

$$\begin{aligned} B_1(x, t_0) &= e^{-r\Delta t} \mathbb{E}^Q[B(A, T)|x] \\ &= e^{-r\Delta t} \int_a^b B(A, T) f(A|x) dA. \end{aligned} \quad (8)$$

In the second step, since $f(A|x)$ is usually not known whereas the characteristic function is, we replace the density by its cosine expansion in A , so we have:

$$f(A|x) = \sum_{k=0}^{\infty} D_k(x) \cos(k\pi \frac{A-a}{b-a}), \quad (9)$$

$$D_k(x) = \frac{2}{b-a} \int_a^b f(A|x) \cos(k\pi \frac{A-a}{b-a}) dA, \quad (10)$$

$$B_1(x, t_0) = e^{-r\Delta t} \int_a^b B(A, T) \sum_{k=0}^{\infty} D_k(x) \cos(k\pi \frac{A-a}{b-a}) dA. \quad (11)$$

We interchange the summation and integration, and insert the definition

$$\mathbb{B}_k = \frac{2}{b-a} \int_a^b B(A, T) \cos(k\pi \frac{A-a}{b-a}) dA. \quad (12)$$

Therefore, we have

$$B(x, t_0) = e^{-r\Delta t} \sum_{k=0}^{N-1} \text{Re}\{\phi(\frac{k\pi}{b-a}, x) e^{-ik\pi \frac{a}{b-a}}\} \mathbb{B}_k, \quad (13)$$

with characteristic function ϕ .

We first explain the recursion procedure for recovering the characteristic function of the

$$R_j = \log(\frac{S_j}{S_{j-1}}), \quad j = 1, \dots, M,$$

a stochastic process, Y_j , is introduced, where $Y_1 = R_M$ and for $j = 2, \dots, M$ we have

$$Y_j := R_{M+1-j} + \log(1 + \exp(Y_{j-1})),$$

and we have that

$$\sum_{j=0}^M S_j = (1 + \exp(Y_M))S_0.$$

Here, however, we will recover the characteristic function of Y_M instead, by a forward recursion procedure, which is then used in turn to recover the transitional density of the European-style arithmetic mean of the underlying process in the risk-neutral formula (14). The CAT bond price is now defined as:

$$B(x_0, t_0) = e^{-r\Delta t} \int_{-\infty}^{+\infty} B(y, T) f_{Y_M}(y) dy. \quad (14)$$

Recovery of characteristic function. We apply the Fourier cosine expansion to approximate $f_{Y_{j-1}}(x)$, giving

$$\hat{\phi}_{Z_{j-1}} = \frac{2}{b-a} \sum_{l=0}^{N-1} \text{Re}(\hat{\phi}_{Y_{j-1}}(\frac{l\pi}{b-a}) \exp(-ia \frac{l\pi}{b-a})) \cdot \int_a^b (e^x + 1)^{iu} \cos((x-a) \frac{l\pi}{b-a}) dx, \quad (15)$$

$\hat{\phi}_{Y_{j-1}}$ is an approximation of $\phi_{Y_{j-1}}$.

$$\Phi_{j-1} = \mathbf{M} D_{j-1},$$

using

$$\Phi_{j-1} = (\Phi_{j-1}(k))_{k=0}^{N-1}, \quad \Phi_{j-1}(k) = \hat{\phi}_{Z_{j-1}}(u_k), \quad u_k = \frac{k\pi}{b-a}, \quad k = 0, \dots, N-1$$

$$\mathbf{M} = (M(k, l))_{k, l=0}^{N-1}, \quad \mathbf{M}(k, l) = \int_a^b (e^x + 1)^{iu_k} \cos((x-a)u_l) dx,$$

$$D_j = \frac{2}{b-a} (D_j(l))_{l=0}^{N-1}, \quad D_j(l) = \text{Re}(\hat{\phi}_{Y_{j-1}}(u_l) \exp(-ia u_l)).$$

Pricing of CAT bond:

$$\hat{B}(x, t_0) = e^{-r\Delta t} \sum_{k=0}^{N-1} \text{Re}(\hat{\phi}_{Y_M}(\frac{k\pi}{b-a}) e^{-ik\pi \frac{a}{b-a}}) \mathbb{B}_k. \quad (16)$$

Clenshaw-Curtis quadrature. In this section, we denote by n_q the number of terms in the Clenshaw-Curtis quadrature (q stands for quadrature). We discuss the efficient computation of matrix $\mathbf{M}$ in (15). An important feature is that matrix $\mathbf{M}$ remains constant for all time steps $t_j, j = 1, \dots, M-1$, so that we need to calculate it only once. Its elements are given by:

$$\mathbf{M}(k, l) = \int_a^b (e^x + 1)^{iu_k} \cos((x-a)u_l) dx, \quad k, l = 0, \dots, N-1, \quad (17)$$

Here (17) is approximated numerically by the Clenshaw-Curtis quadrature rule, which is based on an expansion of the integrand in terms of Chebyshev polynomials (as proposed in [1]). The Clenshaw-Curtis as well as the Gaussian quadrature rules exhibit an exponential convergence for the integration in (17), but the Clenshaw-Curtis quadrature is preferred here, since it is computationally cheaper. In [1], it is shown when using the Clenshaw-Curtis quadrature rule to compute matrix $\mathbf{M}$ (only once, used for all time steps), the total computational complexity is thus $O(n_q \log_2 n_q) + O(n_q N^2)$. Furthermore, at each time step t_j , we need $O(N^2)$ computations for the matrix-vector multiplication (15) and $O(N)$ computations to obtain $\hat{\phi}_{Y_j}$. The computational complexity for this task is thus $O(MN^2)$. The overall computational complexity of

pricing method for CAT bond is then $O(n_q \log_2 n_q) + O(n_q N^2) + O(MN^2)$. The number N^2 is in practice much larger than $\log_2 n_q$. The overall complexity is then of order $O((n_q + M)N^2)$. In [7], it is proved that for most exponential Levy processes, the Fourier cosine expansion exhibits an exponential convergence rate with respect to N . For the integrand in (17) the Clenshaw-Curtis quadrature converges exponentially with respect to n_q . Therefore, the pricing method is an efficient alternative to the method proposed in [3], which requires $O(M\bar{N}^2)$ computations ($\bar{N}$ being the number of points used in the quadrature in [3]), with $\bar{N} > n_q$, as well as $\bar{N} > N$, for the same level of accuracy.

References

- [1] C. W. Clenshaw and A. R. Curtis, A method for numerical integration on an automatic computer, *Numer. Math.*, pp. 197205, (1960).
- [2] F. Fang and C. W. Oosterlee, A novel pricing method for European options based on Fourier-cosine series expansions, *SIAM J. Sci. Comput.*, 31, pp. 826848, (2008).
- [3] G. Fusai and A. Meucci, Pricing discretely monitored Asian options under Levy processes, *J. Bank. Finance*, pp. 20762088, (2008).
- [4] D. Lemmens, L. Liang, J. Tempere, and A. De Schepper, Pricing bounds for discrete arithmetic Asian options under Levy models, *Phys. A*, 389, pp. 51935207, (2010).
- [5] A.J.A. Unger, Pricing index-based catastrophe bonds, part1: formulation and discretization issues using a numerical PDE approach. *Computers and Geosciences*, inpress, doi:10.1016/j.cageo.2009.06.010, (2009).
- [6] A.J.A. Unger, Pricing index-based catastrophe bonds: Part 2: Object-oriented design issues and sensitivity analysis. *Comput. Geosci.* 36(2), 150160, (2010).
- [7] B. Zhang and C.W. Oosterlee. Efficient pricing of European-style Asian options under exponential Levy processes based on Fourier cosine expansions. *SIAM J. Finan. Math.*, 4:399-426, (2013).

e-mail: saharyaghoby@gmail.com

e-mail: ranjbar633@gmail.com

e-mail: bastani@iasbs.ac.ir

Stochastic optimization in correlated multiple insurance business lines

Khaled Masoumifard¹

Shahid Beheshti, Tehran, Iran

Mohammad Zokaei

Shahid Beheshti, Tehran, Iran

Abstract

The present paper addresses the issue of the stochastic control of the optimal dynamic reinsurance policy and dynamic dividend strategy, which are state-dependent, for an insurance company that operates under multiple insurance business lines. The aggregate claims model with a thinning-dependence structure is adopted for the risk process. In the optimization method, the maximum of the cumulative expected discounted dividend payouts with respect to the dividend and reinsurance strategies are considered as value function. This value function is characterized as the smallest super Viscosity solution of the associated Hamilton-Jacobi-Bellman (HJB) equation.

Keywords: Thinning dependence; Hamilton-Jacobi-Bellman equation; Viscosity solution; Dynamic programming principle

AMS Mathematical Subject Classification [2018]: 60J25, 91B30

1 Introduction

Suppose a insurance company based on a dynamic strategy distributes a ratio of its dividend amongst the shareholders and transfers a part of its risk to a secondary insurance company by a dynamic reinsurance strategy. The dividend and reinsurance strategies are shown as $\{D_t\}_{t \geq 0}$ and $\{R_t\}_{t \geq 0}$, respectively. A paramount issue for an insurance company is the optimization of these strategies. For this reason, first an objective function should be considered and then $\{D_t\}_{t \geq 0}$ and $\{R_t\}_{t \geq 0}$ strategies should be found as such that the objective function is optimized. A very common function in literature is the cumulative expected discounted dividends which is displayed as $V(\cdot)$. In the following, we will outline some research on thinning-dependence structure and optimization $V(\cdot)$ with respect to the dividend and reinsurance. Some studies have addressed this issue(e.g. [2] and [3]).

¹speaker

2 Problem formation

A control strategy is a process $\pi = (\mathbf{R}, D)$ where $\mathbf{R}$ is a vector of reinsurance strategies and D_t is a dividend strategy. Reinsurance can be an effective way to manage risk by transferring risk from an insurer to a second insurer (referred to as the reinsurer). A reinsurance contract is an agreement between an insurer and a reinsurer under which, claims that arise are shared between the insurer and reinsurer. Let a Borel measurable function $R : [0, \infty) \rightarrow [0, \infty)$, be called retained loss function, describing the part of the claim that the company pays and satisfies $0 \leq R(\alpha) \leq \alpha$. The reinsurance company covers $\alpha - R(\alpha)$, where the size of the claim is α . Now assume that in order to reduce the risk exposure of the portfolio, the insurer can take reinsurances in a dynamic way for some insurance lines, each of these reinsurances is indexed by $\{1, \dots, n\}$. We denote by $\mathcal{F}$ the vector $(\mathcal{F}_1, \dots, \mathcal{F}_n)$, in which $\mathcal{F}_i$ is the family of retained loss functions associated to the reinsurance policy in i 'th line. Thus, the reinsurances control strategy is a collection $\mathbf{R} = (\mathbf{R}_t)_{t \geq 0} = (R_{1t}, \dots, R_{nt})_{t \geq 0}$ of the vector functions $\mathbf{R}_t : \Omega \rightarrow \mathcal{F}$ for any $t \geq 0$.

Well-known reinsurance types are:

- (1) Proportional reinsurance with $R_P(\alpha) = b\alpha$,
- (2) Excess of loss reinsurance (XL) with $R_{XL}(\alpha, M) = \min\{\alpha, M\}$, $0 \leq M \leq \infty$.
- (3) Limited XL reinsurance (LXL) with $R_{LXL}(\alpha, M) = \min\{\alpha, M\} + (\alpha - M - L)^+$, $0 \leq M, L \leq \infty$.

The numbers M and L are named priority and limit, respectively.

A dividend strategy is a process $D = (D_t)_{t \geq 0}$ where D_t is the cumulative amount of dividends paid out by the reinsurance. Denote by Π_x the set of all control strategies with initial surplus $x \geq 0$. Now, for any $\pi \in \Pi_x$, the surplus process can be written as

$$X^\pi(t) \stackrel{st}{=} x + \int_0^t p_{\mathbf{R}_s} ds - \sum_{i=1}^{N'_t} Z_i - D_t \quad (1)$$

where N'_t is a Poisson process with calim arrival intensity $\beta = \sum_{i=1}^m \beta_i$ and the Z_i are i.i.d random variable with distribution

$$G^{\mathbf{R}}(\alpha) = \sum_{j=1}^n \sum_{k=1}^{\binom{n}{j}} \left[\sum_{i=1}^m \frac{\beta_i}{\sum_{i=1}^m \beta_i} \prod_{z \in A_{jk}^n} p_{iz} \prod_{z \in S^n - A_{jk}^n} (1 - p_{iz}) \right] F_{\mathbf{R}_{A_{jk}^n}}(\alpha) \quad (2)$$

where $F_{\mathbf{R}_{A_{jk}^n}}(\alpha) = p(\sum_{z \in A_{jk}^n} R_z(U_z) \leq \alpha)$. The time of ruin for this process is defined by

$$\tau^\pi = \inf \{t \geq 0 : X^\pi(t) < 0\}. \quad (3)$$

In this paper, we assume that the reinsurance calculates its premium using the expected value principle with reinsurance safety loading factor $\eta_1 \geq \eta > 0$:

$$q_R = (1 + \eta_1) \left(\sum_{i=1}^m \beta_i \right) E(Y - Z) = (1 + \eta_1) \left(\sum_{i=1}^m \beta_i \right) \left(\int_0^\infty \alpha dG(\alpha) - \int_0^\infty \alpha dG^{\mathbf{R}}(\alpha) \right)$$

and so $p_R = p - q_R$, where

$$p = (1 + \eta) \left(\sum_{i=1}^m \beta_i \right) E(Y) = (1 + \eta) \left(\sum_{i=1}^m \beta_i \right) \int_0^\infty \alpha dG(\alpha).$$

A limitation of the existing in this model is the implicit or explicit assumption that the insurers produce only one type of insurance, even though most insurers produce multiple types of coverage (e.g., automobile insurance, general liability insurance, fire insurance, workers' compensation insurance, etc.). The dependency can be introduced between the processes through thinning: suppose that an insurance company has n ($n \geq 2$) lines of business and stochastic sources that may cause a claim in at least one of the n lines are classified into m class. It is assumed that each event in the k th class may cause a claim in the j th line with probability p_{kj} for $k = 1, 2, \dots, m$ and $j = 1, 2, \dots, n$. Regarding an admissible control strategy $(\pi_t)_{t \geq 0}$ and an initial reserve $x \geq 0$, we define the following value function:

$$V^\pi(x) = E_x \left[\int_0^{\tau^\pi} e^{-\delta s} dD_s \right].$$

Our aim in this paper is to extend this result for the model described earlier, in other words, we are looking for

$$V(x) = \sup_{\pi \in \Pi_x} V^\pi(x). \quad (4)$$

To obtain the Hamilton-Jacobi-Bellman (HJB) equation associated with the value function (4), we need to state the so-called Dynamic Programming Principle (DPP). So, the HJB equation can be written as

$$\max\{1 - V'(x), \sup_{\mathcal{F}} \mathcal{L}_{\mathbf{R}}(V)(x)\} = 0. \quad (5)$$

where

$$\mathcal{L}_{\mathbf{R}}(V)(x) = p_{\mathbf{R}} V'(x) - (\delta + \sum_{i=1}^m \beta_i) V(x) + (\sum_{i=1}^m \beta_i) \int_0^x V(x - \alpha) dG^{\mathbf{R}}(\alpha). \quad (6)$$

2.1 Dividend band strategy with reinsurance

Let $\mathcal{A}$, $\mathcal{B}$, and $\mathcal{C}$ are disjoint sets with $\mathcal{A} \cup \mathcal{B} \cup \mathcal{C} = \mathbb{R}_+$, we say $\mathcal{P} = (\mathcal{A}, \mathcal{B}, \mathcal{C})$ is a band partition if $\mathcal{A}$ is closed, bounded, and nonempty; $\mathcal{C}$ is open from the right; $\mathcal{B}$ is open from the left, the lower limit of any connected component of $\mathcal{B}$ belongs to $\mathcal{A}$, and there exists $b \geq 0$ such that $(b, \infty) \in \mathcal{B}$.

Definition 2.1. Consider an initial surplus $x \geq 0$, a stationary reinsurance control $\mathbf{r}^x = (r_1^x, \dots, r_n^x)$, and a band partition $\mathcal{P} = (\mathcal{A}, \mathcal{B}, \mathcal{C})$. An admissible control strategy $\pi^x = (\mathbf{R}^x, D^x) = (\mathbf{R}_t^x, D_t^x)_{t \geq 0} \in \Pi_x^\pi$ is define as follows,

- if $x \in \mathcal{A}$, we set $D_t^x = p_{\mathbf{r}} t = \sum_{i=1}^n p_{r_i^x} t$ and $\mathbf{R}_t^x = \mathbf{r}^x$. Afterward, follow the strategy corresponding to initial surplus $x - \mathbf{r}^x(U_1)$ where U_1 is the size of first claim and $\mathbf{r}^x(U_1) = \sum_{i=1}^n r_i^x(U_1) I_{U_1 \in L_i}$, where $U_1 \in L_i$ indicates that U_1 is a claim from the line i .
- if $x \in \mathcal{B}$, there exists $x_0 \in \mathcal{A}$ such that $(x_0, x) \subset \mathcal{B}$, then we set $L_0^x = x - x_0$ and $\mathbf{R}_0^x = \mathbf{r}^x$. Afterward, follow the strategy corresponding initial surplus x_0 .
- if $x \in \mathcal{C}$, there exists $x_1 \in \mathcal{A}$ such that $(x_1, x) \subset \mathcal{B}$. Then $D_t^x = 0$ and $\mathbf{R}_t^x = \mathbf{r}^{X_{t-}}$ up to $\tau' = \inf\{t : X_t \notin \mathcal{C}\}$. Afterward, follow the strategy corresponding initial surplus $X_{\tau'}$.

The family $\pi(\mathcal{P}, \mathbf{r}) = \{(\mathbf{R}^x, D^x) \in \Pi_x^\pi, x \geq 0\}$ is called the reinsurance band strategy associated with $\mathcal{P}$ and $\mathbf{r}$.

3 Main results

In this section, we state a comparison result between viscosity subsolutions and supersolutions of (5) with a suitable boundary condition that gives us the uniqueness of viscosity solution. Also, we characterize the optimal value function as the smallest supersolution of the HJB equation.

Definition 3.1. We say that a function $u : [0, \infty) \rightarrow \mathbb{R}$ belongs class L if satisfies

- (i) u is locally Lipschitz,
- (ii) if $0 \leq x < y$, then $u(y) - u(x) \geq y - x$, and
- (iii) there exists a constant $k > 0$ such that $u(x) \leq x + k$ for all $x \in [0, \infty)$.

We also define; $L^* = \{ u : u \text{ is viscosity solution of (5) and belongs to } L \}$.

It is interesting to note that if u is of class L , then u is strictly positive, linearly bounded, nondecreasing and absolutely continuous.

Proposition 3.2. *The optimal value function $V(x) = \sup_{\pi \in \Pi_x} V^\pi(x)$ is the smallest viscosity supersolution of (5) that belongs to L .*

These results allows us to characterize V as the unique viscosity solution of (5) with boundary condition $V(0) = \inf_{u \in L^*} u(0)$. From the previous proposition we can deduce the usual viscosity verification result: If we can find a stationary reinsurance strategy $\pi = (\mathbf{R}^x, D^x) \in \Pi_x$ such that V^π is a viscosity supersolution of (5), then $V(x) = V^\pi(x)$; because $V(x) \geq V^\pi(x)$ and by above proposition $V(x)$ is the smallest viscosity supersolution of (5). Now we can show that the optimal control strategy is a reinsurance band strategy.

Theorem 3.3. *Let the vector $\mathcal{F} = (\mathcal{F}_1, \dots, \mathcal{F}_n)$, where $\mathcal{F}_i$ is one of the reinsurance families; proportional reinsurance family ($\mathcal{F}_P$), excess of loss reinsurance family ($\mathcal{F}_{XL}$) and limited excess of loss reinsurance family ($\mathcal{F}_{LXL}$). Then, there exists an admissible reinsurance control $\mathbf{R}^* \in \mathcal{F}$ such that $\pi(\mathcal{P}^*, \mathbf{R}^*)$, the reinsurance band strategy associated to $\mathcal{P}^*$ and $\mathbf{R}^*$, is optimal.*

4 Numerical results

Now, we obtain numerically some examples by using the above algorithm.

Example 4.1. Let insurance company has three lines of business such that it's risk process has the Thinning-dependence structure, $F_i(x) = 1 - e^{-\lambda_i x}$, $i = 1, 2, 3$, and $\lambda_1 = 0.5$, $\lambda_2 = 3$, $\lambda_3 = 2$, $\beta_1 = 8$, $\beta_2 = 4$, $\beta_3 = 5$, $\eta = 3$, $\eta_1 = 3.5$, $p_{11} = 1$, $p_{12} = 0.06$, $p_{13} = 0.05$, $p_{21} = 0.03$, $p_{22} = 1$, $p_{23} = 0.01$, $p_{31} = 0.007$, $p_{32} = 0.005$, $p_{33} = 1.0$ and $\delta = 0.3$. The reinsurance strategy in i th line is depicted by R_i . As was mentioned before, if $R_i \in \mathcal{F}_P$ then $R_i(y) = b_i(\cdot)y$ and if $R_i \in \mathcal{F}_{XL}$ then $R_i(y) = \min(y, M_i(\cdot))$, where $b_i(\cdot)$ and $M_i(\cdot)$ are functions of the company's capital. If the insurance company considers a reinsurance contract for three lines, the optimization issue will be equal with the uni-dimensional model scrutinized by [?]. Using the recently explained numerical method, the following results are gleaned,

- (i) if $R_i \in \mathcal{F}_P$, $i = 1, 2, 3$, then, $\mathcal{P} = (\{12.26\}, (12.26, \infty), [0, 12.26])$,

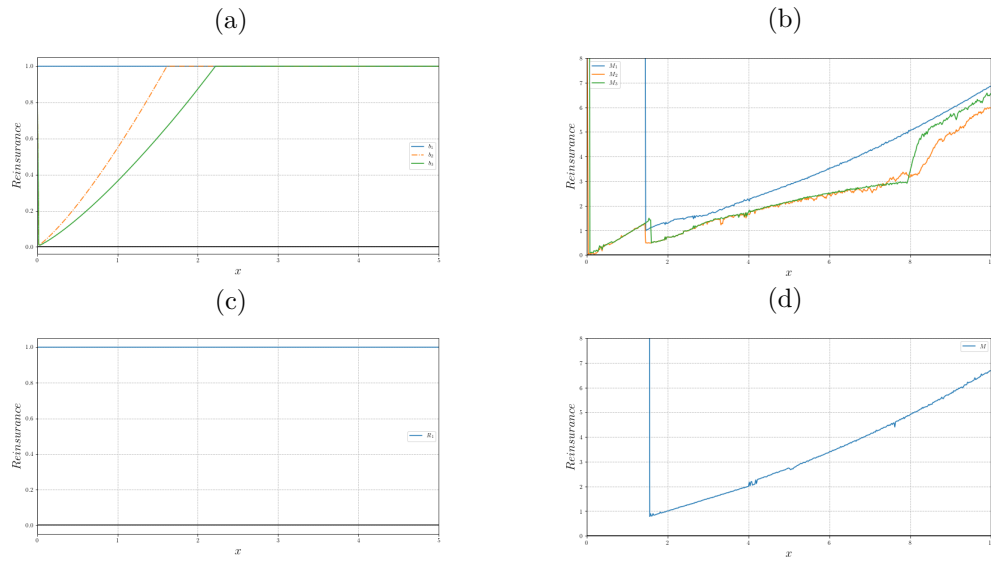


Figure 1: The numerical solution of the optimal reinsurances with $h = 0.02$, (a) The optimal results when three proportional reinsurances are used for three lines, (b) The optimal results when three XL reinsurances are used for three lines, (c) The optimal result when one proportional reinsurances is used for three lines, (d) The optimal results when one XL reinsurance is used for three lines

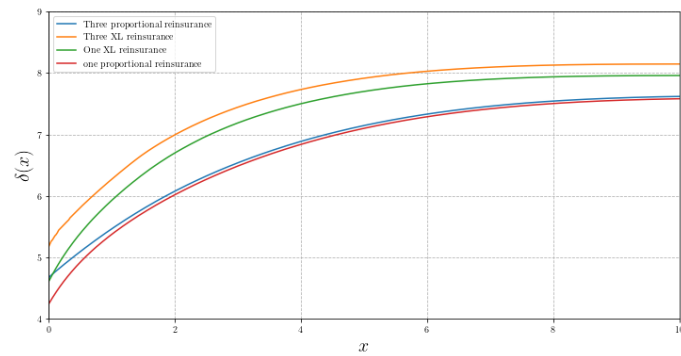


Figure 2: Survival functions

- (ii) if $R_i \in \mathcal{F}_{XL}$, $i = 1, 2, 3$, then, $\mathcal{P} = (\{10.44\}, (10.44, \infty), [0, 10.44))$,
- (iii) if $R_1 = R_2 = R_3 \in \mathcal{F}_p$, then, $\mathcal{P} = (\{12.3\}, (12.3, \infty), [0, 12.3))$,
- (iv) if $R_1 = R_2 = R_3 \in \mathcal{F}_{XL}$, then, $\mathcal{P} = (\{10.64\}, (10.64, \infty), [0, 10.64))$.

Also, optimization results for the value functions and reinsurance strategy are reported in Figures 1 and 2.

References

- [1] B. J. Christopher and C. David, *On optimal reinsurance, dividend and reinvestment strategies*, Economic Modelling, n. 28, (2011), pp. 211-2018.
- [2] B. J. Christopher and C. David, *Optimal dividends under reinsurance*, Centre for Actuarial Studies, Department of Economics, University of Melbourne Melbourne (2007).
- [3] P. Azkue and N. Muler, *Stochastic optimization in insurance: a dynamic programming approach*, Springer, (2014).

e-mail: k_masoumifard@sbu.ac.ir
e-mail: zokaei@sbu.ac.ir

Designing an Automated Trading System Using Image Processing by a Convolutional Neural Network

Amir Hossein Yaftian¹

Tarbiat Modares University, Tehran, Iran

Mohammad Ali Rastegar

Tarbiat Modares University, Tehran, Iran

Abstract

Artificial Intelligence and Machine Learning techniques have always been popular for use in Automated Trading systems. This study aims to Design an Automated Trading system using image processing by a 2D Convolutional Neural Network. To do this, at first 28 technical analysis indicators are selected and the values of each are calculated. These values convert to 2D images, as a result, we have a 28×28-dimension image for each point in the time series of price data. Then, each image is labeled with buy, sell or hold. These data are entered to the convolutional neural network. The results show that in 80% of cases, the return of this method is higher than the Buy & Hold strategy. Also, in terms of standard deviation and maximum Drawdown, performed better.

Keywords: Algorithmic Trading, Artificial Intelligence, Deep Learning, Convolutional Neural Network, Technical Analysis

1 Introduction

Recent years can be called the years of increasing use of algorithmic trading systems in financial markets around the world. Currently, a large number of trades in financial markets are executed by algorithms. The growing presence of algorithmic trading systems and the need for new systems and algorithms with different functions, have increased the demand of various institutions for the feasibility, design and development of automated trading systems in various countries, and recently in Iran. The main advantage of such systems for investors is the speed and accuracy of information analysis and decision making without interfering with emotions.

In literature, different machine learning models have been used to predict future values. Traditional machine learning models are very popular to predict the stock markets. Some research has directly implemented timeseries predictions on financial data, while others have used technical analysis data and fundamental analysis to improve prediction performance. Artificial neural network models, genetic algorithms, fuzzy rule systems, and hybrid models were among the best choices.

In recent years, deep learning-based prediction/classification models started emerging as the best performance achievers in various applications, outperforming classical computational intelligence methods like SVM. However, image processing and vision-based problems dominate the type of applications that these deep learning models outperform the other techniques [1]. Some researchers have used deep learning techniques such as Recurrent Neural Network (RNN), Convolutional Neural Network (CNN), and Long-Short Term Memory (LSTM). But the use of deep neural techniques in financial forecasting models has been very limited.

Gudelek et al. [2] used 28 technical analysis indicators to create images for each day of the price timeseries data of 17 index funds. Then they labeled these images, once with buy and sell labels, and once again with the buy, sell and hold labels. They imported these images to the CNN model to predict the test data labels. They considered the train data from 2000 to 2014, and the results on the test data from 2015 to 2017 showed that the return of this method is higher than the Buy & Hold method. Sezer and Ozbayoglu [3] also used 15 technical analysis indicators on Dow Jones stocks and a number of ETFs to create images. They labeled these images with buy, sell and hold

¹ speaker

labels and entered them to the CNN model. like the research by Goodluck et al. [2], the results of this research showed that in many cases, the return of this method is higher than other methods. In another research, in order to predict buy, sell and hold labels, Sezer and Ozbayoglu [4] used a different approach to creating images than their previous research. As in previous research, they found that their method was more profitable than the Buy and Hold method.

The model presented in this study is designed for trading in the stock market. For this purpose, a number of Tehran Stock Exchange stocks are selected as samples from the entire market. According to the proposed model of this research, to select from all the market stocks, the following items are considered as criteria for stock selection: Number of data (trading days), Lack of price gap in price data history and Distribution in selected industries. This data was provided by Noavarane Amin Financial Information Processing Company.

Table 1: Selected stocks

Thicker	Industry	Start date	End date	Number of data	Average Annual Return	StdDev	Skewness	Kurtosis
Cefars	Cement	1380/01/05	1398/05/30	3866	0.45920	0.02572	0.66276	57.1232
Felooleh	Basic metals	1380/01/06	1398/05/30	3424	0.49078	0.02764	-0.39341	24.8828
Kerooy	Metallic ores	1380/01/20	1398/05/30	3732	0.43775	0.03019	1.93337	42.6031
Khodro	Auto manufacturer	1380/01/05	1398/05/30	3827	0.441308	0.02456	-1.75556	59.1761
Shebehran	Petroleum	1380/01/27	1398/05/30	3691	0.533963	0.01931	1.83715	18.1807

After receiving the data, the label of each data is determined. In this study, the local minimum and local maximum points are determined through an 11-day trailing window. local minimum points receive buy labels and the local maximum points receive sell labels. The rest of the points that are neither local maximum nor local minimum, receive hold labels.

The CNN model receives images as input. To create these images, we use the following technical analysis indicators. These indicators have been selected based on previous researches.

Table 2: Technical analysis indicators and their parameters for creating images

Name	Parameters	Number
$\tanh((\text{second differential}(\text{close})))$	-	1
Volume	-	1
RSI	15-20-25-30	4
SMA	15-20-25-30	4
MACD	26*12 - 28*14 - 30*16	3
MACD trigger	9*26*12 - 10*28*14 - 11*30*16	3
William %R	14-18-22	3
Stochastic Oscillator	14-18-22	3
Ultimate Oscillator	7*14*28 - 8*16*22 - 9*18*36	3
MFI	14-18-22	3

after calculating the above table values for all trading days of selected stocks, we normalize each of these 28 features between -1 and 1.

Now, we can create images. To do this, we create a 28×28 matrix for each day of the stock price. First, we consider the twenty-eighth days ago. We put all 28 features listed in the table above for this day in the first column of the matrix. Now we go one day ahead, the twenty-seventh day, and we put the values of all these 28 features for this day in the second column of the matrix. We go ahead until the 28th column of the matrix that shows the current day and do the same. Now we have a 28×28 matrix.

The approach of this research for splitting data into train and test sets is adopted by Cross Validation method. To this end, we split each timeseries data into a number of batches with 250 Data. For example, we divide the

Khodro time series data into 12 subsets with 250 Data. Now we consider the first five batches as the training set and the next batch as the testing set. In the next step we will move 1 batch forward and repeat the same thing again.

Another important point is the data imbalance. The number of buy and sell labels is much lower than the number of hold labels. So, we need to increase the number of buy and sell labels. There is no specific formula to determining the best ratio of each labels number to total labels number, but generally the closer the ratios are to each other, the better. To increase these ratios, we duplicated the data labeled buy and sell.

Finally, we create the model. In this study, a CNN model was used. A Convolutional Neural Network (CNN) is a Deep Learning algorithm which can take in an input image, assign importance (learnable weights and biases) to various aspects/objects in the image and be able to differentiate one from the other. The pre-processing required in a CNN is much lower as compared to other classification algorithms. While in primitive methods, filters are hand-engineered, with enough training, CNN have the ability to learn these filters/characteristics.

In general, the four main layers that make up a CNN network are: Convolution layer, Pooling layer, Dropout layer and fully connected layer. Each layer has its own different task. Model training in each convolutional neural network, like other neural networks, consists of two steps, feed forward and backpropagation. In the first step, or feed forward, the input images are imported into the network. This is actually is a point multiplication between the input image data and the parameters of each neuron. Next, the convolution operation applies to each layer and then the network output calculates. Now to adjust the network parameters, the error rate is calculated using the result of the computational output. To do this, the output of the network compared with the correct response using an error function. in the next step, according to the calculated error rate, the backpropagation phase begins. During this step, according to the chain rule, the gradient of each parameter is calculated, and then all the parameters are changed according to the effect that each parameter has on the error in the network. Finally, after updating the parameters, the next step of feed-forward begins. After a good number of these steps have been repeated, the network training will end. The proposed model of this research is as follows:

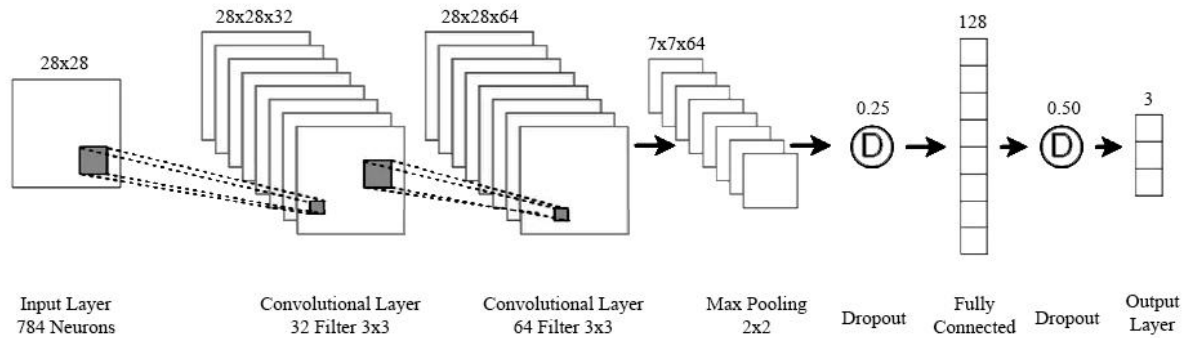


Figure 1: The proposed CNN model

2 Main results

After 75 epochs, account Balance and Recall ratio will decrease. So, we consider 75 epochs as the optimal number of epochs and examine the model and results with this number of epochs.

Table 3: Confusion Matrix and Related ratios for Khodro stock

		Predicted		
		Hold	Buy	Sell
Actual	Hold	1269	135	148
	Buy	18	78	0
	Sell	8	0	94
Recall		0.8176	0.8125	0.9215
Precision		0.9799	0.3661	0.3884
F1 Score		0.8914	0.5048	0.5465

In all stocks, Recall shows a high percentage, but Precision is low. This can be for two reasons. The first reason is over-fitting of the model. As the epochs increases, the model identifies other data with hold labels as buy and sell. Because the number of buy and sell labels in the test set is 10% of the total label numbers. This shows the imbalance in the data. The second reason is the similarity of hold labels points with near points that have buy and sell labels. The value of technical analysis indicators at these points is approximately close to the points with buy and sell labels, and the model classify them as buy and sell.

Table 4: Financial Results

		CNN without Commission	CNN with Commission	Buy & Hold
Average Annual Return	Cefars	0.2331	0.0389	0.1568
	Felooleh	0.2249	0.0396	0.1982
	Kerooy	0.2886	0.0939	0.2559
	Khodro	0.3047	0.1261	0.1621
	Shebehran	0.3108	0.1014	0.4274
StdDev	Cefars	0.0712	0.0702	0.1376
	Felooleh	0.0780	0.0768	0.1775
	Kerooy	0.0834	0.0821	0.1450
	Khodro	0.0962	0.0948	0.1725
	Shebehran	0.0750	0.0739	0.1456
Maximum Drawdown	Cefars	0.3068	0.7911	2.1466
	Felooleh	0.2067	0.6187	1.2003
	Kerooy	0.4722	0.5402	0.5478
	Khodro	0.4711	0.6091	1.2874
	Shebehran	0.2494	0.8407	0.6762

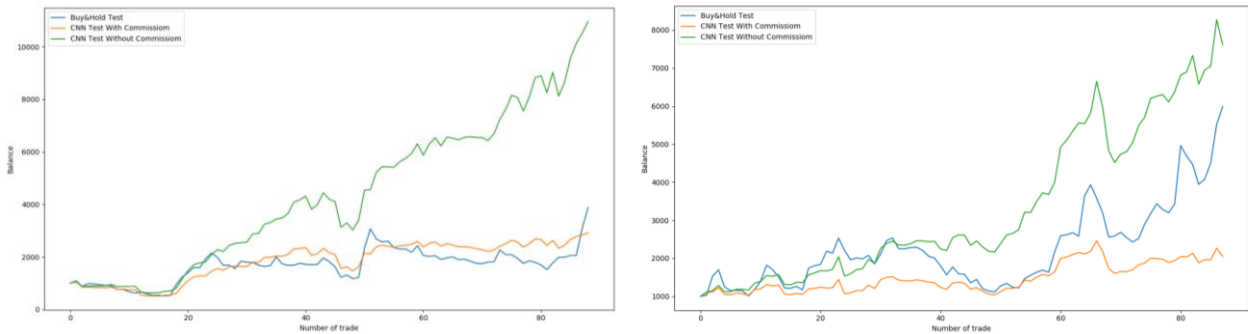


Figure 2: Account balance for Khodro (left) and Kerooy (right) stocks. It can be seen that the account balance without commission (green) is much higher than it with commission (orange). Account balance for Buy and Hold strategy shown in blue.

It can be seen that the return of the CNN model without commission for Cefars, Felooleh, Kerooy and Khodro stocks is higher than other methods. because of the commission, the return of the CNN model with commission in every 5 stocks, is the lowest compared to other methods. The reason of this, is in the average return of each trade. The average return of each trade, for example on Cefars, is 2.3%. Given the 1.4% commission, much of the return is spent on commission. In the case of the risk, the maximum drawdown of the CNN model without commission in all stocks, is lower than other methods.

On the other hand, the results show that the CNN methods have been idle for more than 50% of the test period. That's mean, in some periods there is no stock in the portfolio. This also indicates that the model does not utilize the maximum available resources. Improvements in this area can increase the return of the CNN methods.

References

- [1] Canziani A, Paszke A, Culurciello E. *An analysis of deep neural network models for practical applications*. arXiv preprint arXiv:1605.07678. 2016 May 24.
- [2] Gudelek MU, Boluk SA, Ozbayoglu AM. *A deep learning-based stock trading model with 2-D CNN trend detection*. In 2017 IEEE Symposium Series on Computational Intelligence (SSCI) 2017 Nov 27 (pp. 1-8). IEEE.
- [3] Sezer OB, Ozbayoglu AM. *Algorithmic financial trading with deep convolutional neural networks: Time series to image conversion approach*. Applied Soft Computing. 2018 Sep 1; 70:525-38.
- [4] Sezer OB, Ozbayoglu AM. *Financial Trading Model with Stock Bar Chart Image Time Series with Deep Convolutional Neural Networks*. arXiv preprint arXiv:1903.04610. 2019 Mar 11.

Email: amirhoseinyaftian@modares.ac.ir

Email: ma_rastegar@modares.ac.ir

Model Selection for Value at Risk with Machine Learning Methods

MohammadTaha Hoveizavi¹
Kharazmi University, Tehran, Iran

Hirbod Assa
Institute for financial and actuarial mathematics University of Liverpool, UK

Seyed-Mohammad-Mahdi Kazemi
Department of Financial Mathematics, Faculty of Financial Science, Kharazmi University, Tehran, Iran

Abstract

Risk measurement plays an important role in quantifying risk and risk management. Value at risk is one of the most popular measures that is used in risk management. Various methods have been developed to calculate Value at Risk of a risk capital according to different economic conditions. So it is important to find the best calculation method for different. In this Paper The main idea is to use Machine Learning Methods for forecasting backtesting on the S&P 500 index VaR models.

Keywords: Value at Risk, Machine Learning, Backtesting.

1 Introduction

Risk measurement plays an important role in quantifying risk and risk management. Value at Risk is the maximum amount of loss over a given horizon of time at a certain confidence level[1]. According to the definition of VaR, a variety of methods have been developed for calculating VaR, that can be categorized into three separate groups, parametric methods, semi-parametric methods, and non-parametric methods[2]. Parametric methods represent a group of methods that consider a specific distribution to calculate Value at Risk. By assuming a normal distribution for return data VaR can be simply calculated as follow[2], formula (1):

$$VaR_t = -P_{t-1}(\mu_t - \sigma Z_\alpha). \quad (1)$$

Where here VaR_t is value at risk at Time t , $-P_{t-1}$ is the stock price at time $t - 1$, μ_t is the return average for period t , σ is the standard deviation for the period t and Z_α is the normal standard amount in $1 - \alpha$ percent of confidence level. Hereinafter, called the Value at Risk with the normal distribution at 95% confidence level $NormalVaR_\alpha$.

Semi-parametric methods have been proposed to estimate VaR, such as application of Extreme Value Theory [3], Historical and Monte Carlo simulation. They are two models based on non-parametric methods that do not assume any distribution for data[2]. Various statistical tests have been developed to validate VaR methods, Binomial test(Bin)[4], Proportion of failures test(pof)[5], Conditional coverage mixed test(CC)[6], Conditional coverage independence test(CCI)[6], Time between failures mixed test(TBF)[7], Time between failures independence test(TBFI)[7], and Time until first failure test(TUFF)[8] are Seven of the most important of these tests. These statistical tests are the major backtesting methods and are based on last failures. Figure 3.

¹speaker

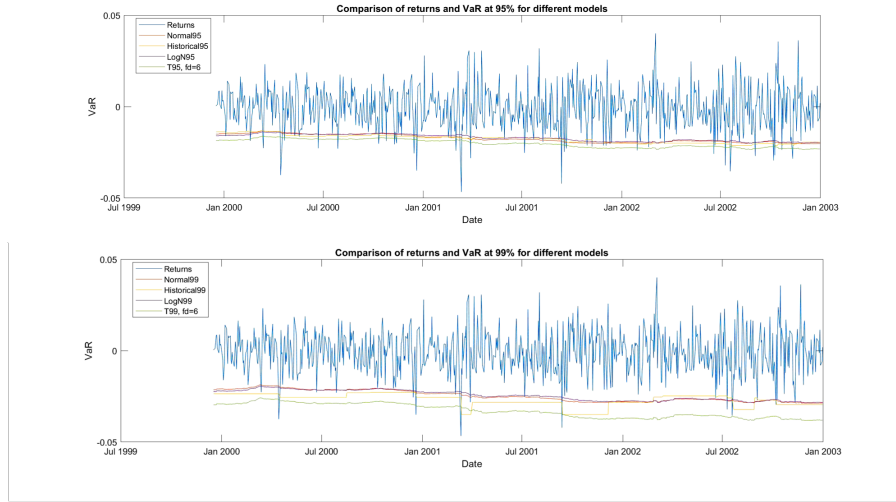


Figure 1: Value at Risk Parametric methods calculated for S&P500 Index from Jan 2000 to Jan 2003.

Machine learning is a statistical tool endowed with computer programming techniques to optimize a performance criterion using example data or past experience[9]. With the advancement of science, the usage of machine learning techniques has increased with the goals of data processing and pattern recognition. There are a lot of problems in finance that have been solved with machine learning techniques, Robert Culkin Sanjiv R. Das (2017) have been used machine Learning to solve high dimensional Black and Scholes (1973) option pricing formula[10], Chongda Liu, Jihua Wang, Di Xiao, Qi Liang (2016) have been tried to forecast S&P500 Movement, they achieved 63 percent accuracy[11]. In this paper, we aim to select the best VaR method in various economic situations and different times. As we said there are some statistical tests that are known as backtesting methods where we can use them to evaluate VaR performance, in Figure2 we can see the results of the backtesting method for $NormalVaR_{95}$. Our idea to choose the most appropriate VaR method is to use Machine Learning Classification to predict backtesting Tests for Value at Risk methods in the next 5 days horizon. In this prediction, the method that has been succeeded to be accepted from each of these seven tests is selected as the appropriate method for calculating VaR on the desired day.

5 Bin	6 POF	7 TUFF	8 CC	9 CCI	10 TBF	11 TBFI
reject	reject	accept	reject	accept	reject	reject
reject	reject	accept	reject	accept	reject	reject
reject	reject	accept	reject	accept	reject	reject
reject	reject	reject	reject	accept	reject	reject
reject	reject	accept	reject	accept	reject	reject
reject	reject	accept	reject	accept	reject	reject
reject	reject	accept	reject	accept	reject	reject
reject	reject	accept	reject	accept	reject	reject
reject	reject	accept	reject	accept	reject	reject
reject	reject	accept	reject	accept	reject	reject
reject	reject	accept	reject	accept	reject	reject

Figure 2: Backtesting results on S&P500 $NormalVaR_{95}$ for eleven consecutive days using a 250-day moving window.

2 Main results

In this paper, we have tried to find the best VaR method at different times via forecasting the Backtesting tests with Machine Learning methods. So we have to assess the performance of the classification methods.

The data used are related to the SP500 index. Features are the p-value of Backtesting methods. The target is zero when the VaR method is accepted and is one if the Backtesting method rejects VaR method. Table 1 shows the quality measures for the test set for each classification technique. We use classification

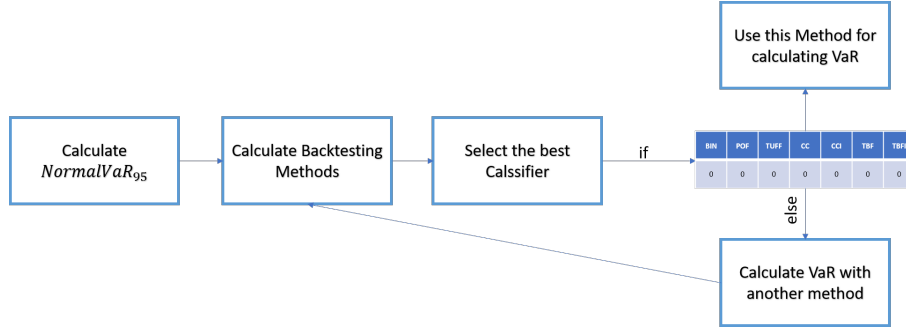


Figure 3: General flowchart of the proposed methodology.

for forecasting POF test with setting 75% of data as train set and 25% of data as test set. Metrics such as Accuracy, Precision, Recall, and F1-Score are key factors for choosing an algorithm, therefore we use the Multiplication of these factors to select the best classifier.

Table 1: Train results for prediction with some of machine learning methods, in this table we set Normal-VaR95's POF test as the target

Classification Method	Train	Test	TN	FN	TP	FP	Acc	F1	precision	Recall
KNN	187	63	59	0	1	3	0.9523	0.4	0.25	1
SVM	187	63	59	0	0	4	0.9365	0	0	0
Gaussian naive bayes	187	63	55	4	4	0	0.9365	0.66	1	0.5
Decision Tree	187	63	59	0	4	0	1	1	1	1
Logistic regression	187	63	59	0	0	4	0.9365	0	0	0
Random Forest	187	63	59	0	3	1	0.9841	0.8571	0.75	1
AdaBoost	187	63	59	0	4	0	1	1	1	1
Balanced Random Forest	187	63	52	7	4	0	0.8888	0.5333	1	0.3636

Table 2: Metrics results for predicting Backtesting tests with the best classifier selected in Table 1

Target	Train	Test	TN	FN	TP	FP	Acc	F1	precision	Recall
Bin	187	63	34	0	29	0	1	1	1	1
POF	187	63	59	0	4	0	1	1	1	1
TUFF	187	63	59	0	3	1	0.9841	0.8571	1	0.75
CC	187	63	38	0	25	0	1	1	1	1
CCI	187	63	35	0	28	0	1	1	1	1
TBF	187	63	42	3	18	0	0.9523	0.9655	0.9333	1
TBFI	187	63	39	1	23	0	0.9841	0.9873	0.975	1

After picking the best classifier for each $NormalVaR_{95}$ target in Table 2 we can predict these targets for the 7 coming days. If the predicted Backtesting methods accepted the $NormalVaR_{95}$ our machine will offer to use $NoemalVaR_{95}$ for calculation Value at Risk. Otherwise, the machine will run this algorithm on another VaR calculation method. we have used a moving window for training the machine from historical data to choosing the best VaR calculation method, utilizing a moving window and regarding the methodology suggested in Figure 3 we calculated VaR for 28 days And we've only been violated (we define violation as the time when the return is less than the Value at Risk) for 4 days this means that we achieved approximately 85% accuracy.

According to the importance of risk management and the vast usage of VaR methods to quantifying Risk has made it important to select the best VaR method, the methodology proposed in this paper improved this cited selection problem.

References

- [1] Choudhry, M., *Advance Praise for The VaR Implementation Handbook*.
- [2] Manganelli, S. and Engle, R.F., 2001. *Value at risk models in finance*.
- [3] McNeil, A.J., 1999. *Extreme value theory for risk managers*. Departement Mathematik ETH Zentrum.
- [4] Jorion, P. Financial Risk Manager Handbook. 6th Edition. Wiley Finance, 2011.
- [5] Kupiec, P. Techniques for Verifying the Accuracy of Risk Management Models. Journal of Derivatives. Vol. 3, 1995, pp. 73–84.
- [6] Christoffersen, P. Evaluating Interval Forecasts. International Economic Review. Vol. 39, 1998, pp. 841–862.
- [7] Haas, M. New Methods in Backtesting. Financial Engineering, Research Center Caesar, Bonn, 2001.
- [8] Kupiec, P. Techniques for Verifying the Accuracy of Risk Management Models. Journal of Derivatives. Vol. 3, 1995, pp. 73–84.
- [9] Parsons, S., 2005. Introduction to Machine Learning by Ethem Alpaydin, MIT Press, 0-262-01211-1, 400 pp. The Knowledge Engineering Review, 20(4), pp.432-433.
- [10] Culkin, R. and Das, S.R., 2017. Machine learning in finance: the case of deep learning for option pricing. Journal of Investment Management, 15(4), pp.92-100.
- [11] Liu, C., Wang, J., Xiao, D. and Liang, Q., 2016. Forecasting sp 500 stock index using statistical learning models. Open journal of statistics, 6(06), p.1067.

Email: taha.ho19@hotmail.com

Email: assa.hirbod@gmail.com

Email: smm.kazemi@khu.ac.ir; mekazemi2000@gmail.com