

Mestrado em Estatística e Gestão da Informação

Master Program in Statistics and Information Management

Estimation of Technical Reserves using Stochastic Methods

⟨Analysis and Risk Management⟩

Bruna Raquel Sebastião Gonçalves

Dissertation submitted in partial fulfillment
of the requirements for the degree of

Master of Science in Master in ⟨Statistics and Information
Management⟩

NOVA Information Management School
Instituto Superior de Estatística e Gestão da Informação

Universidade Nova de Lisboa

Estimation of Technical Reserves using Stochastic Methods

⟨Analysis and Risk Management⟩

Bruna Raquel Sebastião Gonçalves

Dissertation submitted in partial fulfillment
of the requirements for the degree of

Master of Science in Master in ⟨Statistics and Information
Management⟩

Estimation of Technical Reserves using Stochastic Methods

Copyright © Bruna Raquel Sebastião Gonçalves, Information Management School, NOVA University Lisbon.

The Information Management School and the NOVA University Lisbon have the right, perpetual and without geographical boundaries, to file and publish this dissertation through printed copies reproduced on paper or on digital form, or by any other means known or that may be invented, and to disseminate through scientific repositories and admit its copying and distribution for non-commercial, educational or research purposes, as long as credit is given to the author and editor.

ABSTRACT

Insurance companies need to have proper technical reserves in order to stay solvent, due to the liabilities that they take responsibility for. The notion of what are proper reserves is a subject studied by insurers and academics.

For the non-life insurance the claim reserves, amongst other types of reserves, have great significance. Hence the existence of many kinds of reserving methods, either stochastic or deterministic.

This work aims to compare the estimates of claim reserves using the Thomas Mack and the Bootstrap methods for the workers compensation insurance.

Keywords: Claims reserving; Chain Ladder; Thomas Mack; Bootstrap.

RESUMO

Companhias de seguros precisam de ter provisões técnicas apropriadas de forma a se manterem solventes, dados os riscos que estas tomam responsabilidade. A noção de o que é uma provisão adequada é um tema estudado pelas seguradoras.

Para o ramo Não-Vida a provisões de sinistros, de entre outros tipos de provisões, têm grande significância. Consequentemente, existem vários métodos para a sua determinação, quer estocásticos, quer determinísticos.

O propósito deste trabalho é a comparação das estimativas das provisões de sinistros no ramo de acidentes de trabalho usando os métodos de Thomas Mack e de Bootstrap.

Palavras-chave: Provisões para Sinistros; Chain Ladder; Thomas Mack; Bootstrap.

CONTENTS

List of Figures	xiii
List of Tables	xv
1 Introduction	1
2 Basic Concepts	3
2.1 The Claims Process	3
2.2 Data	4
2.3 Estimation Methods	5
2.4 Confidence Intervals Estimation	6
3 Chain Ladder Method/ Thomas Mack	9
3.1 Model Assumptions	10
3.1.1 First Assumption	10
3.1.2 Second Assumption	11
3.1.3 Third Assumption	11
3.2 Mean Squared Error and Standard Error	12
3.3 Verification of Assumptions	13
3.3.1 First Assumption	13
3.3.2 Second Assumption	14
3.3.3 Third Assumption	16
3.4 Practical Application	17
4 Bootstrap Method	23
4.1 Bootstrap Method	24
4.2 Pratical Application	26
5 Conclusion	31
Bibliography	33

LIST OF FIGURES

2.1	Claims Process through time	3
2.2	Run-Off Triangle	4
3.1	Incremental paid claims	17
3.2	Cumulative paid claims	18
3.3	Estimation of Development Factors	18
3.4	Data adjustment to f_j	18
3.5	Individual development factors	19
3.6	Individual development factors ranked	19
3.7	Individual development factors	19
3.8	Individual development factors ranked from smallest to largest	20
3.9	Elements of the L and S sets	20
3.10	Weighted Residuals	21
3.11	Full Triangle with estimated future claim amounts	22
3.12	Estimated reserve by Accident Year	22
3.13	Estimation of $\sigma_j^2, 1 \leq j \leq 8$	22
3.14	Estimation of $\widehat{MSE}(\hat{R}_i)$	22
3.15	Estimation of confidence interval for R_i	22
4.1	Bootstrap Method	24
4.2	Incremental paid claims	26
4.3	Cumulative paid claims	26
4.4	Estimation of Development Factors	26
4.5	Fitted Payments	27
4.6	Unadjusted Person Residuals	27
4.7	Residuals of each Development Year	27
4.8	Residuals of each Occurrence Year	28
4.9	Example of simulated residuals	28
4.10	Example of a set of pseudo-data	28
4.11	Reserves for each occurrence year from the example of the pseudo-data	29

LIST OF FIGURES

4.12 Approximate probability distribution of the Total Reserve	29
4.13 Estimation of confidence interval for R_i	30

LIST OF TABLES

3.1	Results of Z_j and Z variables	21
4.1	Bootstrap Results	29
5.1	Thomas Mack vs. Bootstrap	31

INTRODUCTION

Insurance companies have to deal with risk management and one of the key decisions is to build and maintain proper technical provisions so that they can assure the fulfilment of future responsibilities towards policyholders and third parties.

In Non-Life insurance, the most important technical provisions are the ones associated with claim amounts. This is due to the cost of some claims which linger to get settled and others that were not immediately reported.

The continuous effort to improve the accuracy of the estimation of claims is a worthy task when it comes to the profitability of the insurance industry, and a needed one in order to ensure its solvency, Straub and Grubbs, 1998.

The problem that insurance companies face is how to estimate the proper claim reserves so that whenever these are not enough it can jeopardize the companies' solvency and might arise financial problems. If the reserves are substantial, the profitability might be influenced in some sensitive way, so, it is important to have the right amount of provisions.

Thus, the main questions that come across insurance companies are the following:

- What might be the adequate amount in the claim reserves?
- How are these reserves estimated?

To get these answers, several stochastic methods have been constructed with strong theoretical background making possible to not only estimate reserves but also to obtain error measurements and build confidence intervals around these estimations. Therefore, insurance companies may use different methods in order to better estimate claim reserves.

For our work, we will be using the two different stochastic methods that are as follows:

- Chain Ladder/Thomas Mack method which is the most simple and famous distribution-free method;

- Bootstrap Method that is a simulation algorithm based method to create pseudo-data through re-sampling with replacement;

The dataset used for this analysis was extracted from the ASF statistical archive and it will be used the *R* software to construct each model whilst setting out the assumptions.

Thereupon, we will have as Chapter 2 the introduction of some concepts, such as some general aspects of claim reserves, how do claims process through time, how the data is presented and some notations that will be used.

In Chapter 3, there will be a theoretical explanation of the Chain Ladder / Thomas Mack method and its assumptions followed by a practical application and verification of the assumptions to our dataset.

In Chapter 4, the Bootstrap method will be explained and then applied to the dataset.

BASIC CONCEPTS

2.1 The Claims Process

A Non-Life insurance policy is defined as a contract between two parties, the insurer and the insured where, on the occurrence of special events the insurer, normally an insurance company, has to pay a given amount of money, once or several times, to the insured. These amounts depend on the peculiar circumstances of those events, Taylor, 2012.

The following figure represents the claims process since its occurrence to its closure.

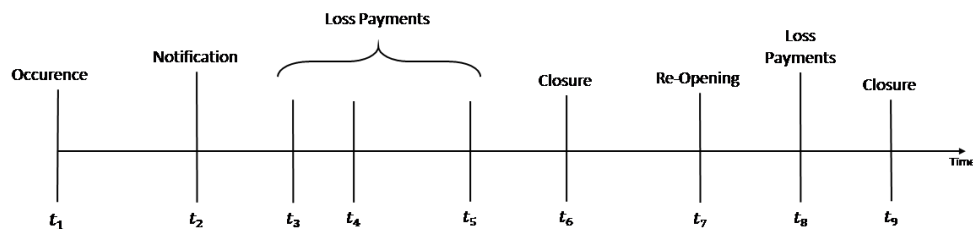


Figure 2.1: Claims Process through time

At time t_1 there is an event that generates a claim and it is called the date of occurrence. Later in time, t_2 , the insured notifies the insurer of the claim, Taylor, 2012.

It is not expected for the insurer to pay the insured immediately. The payments in times t_3 , t_4 and t_5 are made after some necessary proceedings, namely the validation and assessment of the claim. When the insurer consider the end of claim payments, the action on the claim is complete and the file is closed at time t_6 .

Eventually, the need to open the file may arise, t_7 . This can happen due to appearance of new relevant informations related to the claim, and, if it is necessary, the

insurer might have to pay again to the insured, t_8 , in order to rectify the case, and then the file is closed again, t_9 . Note that these $t_k, k = 7, 8, 9$, may occur during several years.

The insurer has the obligation to build provisions, at the end of each practice, that aim to cope with the payments of claims that already happened and weren't yet reported to the insurer, and there might be some claim files that are not yet closed and the payments are due in the next accountants year.

In the interval between t_1 and t_2 , the claim is said to be *incurred but not reported* (IBNR), since the event that led to a claim already happened but the insurer does not know but the latter still takes responsibility for the liability, Taylor, 2012.

When the insurer is already notified that there has been a claim but does not know the real amount that is going to pay to the insured, $[t_2, t_6]$, the provision for this sum is called *incurred but not enough reserved* (IBNER).

2.2 Data

The ASF (Autoridade de Supervisão de Seguros e Fundos de Pensões) is the Portuguese national authority responsible for the conservative and behavioural supervision of the insurance, reinsurance and pension funds activities and its managing entities, “Autoridade de Supervisão de Seguros e Fundos de Pensão”, n.d. In regards of this work, we have retrieved the paid amounts of claims of accidents at work insurances, in thousand of Euros from 2007 until 2016.

To estimate the provision amounts, insurance companies should own historical data regarding claim costs that, through statistical procedures, could deduce its development for the future.

This data is organized and forms a two-dimensional array whose shape is, normally, triangular as we can see in Figure 2.2, Taylor, 2012.

Year of accident	Development Year								
	0	1	2	...	j	...	$n-1$	n	∞
0	$X_{0,0}$	$X_{0,1}$	$X_{0,2}$	...	$X_{0,j}$	...	$X_{0,n-1}$	$X_{0,n}$	$X_{0,\infty}$
1	$X_{1,0}$	$X_{1,1}$	$X_{1,2}$	...	$X_{1,j}$	...	$X_{1,n-1}$		
2	$X_{2,0}$	$X_{2,1}$	$X_{2,2}$	...	$X_{2,j}$	...			
...	...	...	...	...	...				
i	$X_{i,0}$	$X_{i,1}$	$X_{i,2}$	...					
...	...	...	...						
$n-1$	$X_{n-1,0}$	$X_{n-1,1}$							
n	$X_{n,0}$								

Figure 2.2: Run-Off Triangle

Each row of the matrix represents the year of accident, or when that the insurer was notified. The columns of the matrix represents the develop year of the claims, i.e. the number of years that take place since the occurrence of the claims until its settled

by the insurer. The last column, ∞ , is called *ultimate* and holds the information of all occurred claims, but whose settlement will only be done after n development years.

Even though the most common period of time is years, the data can be represented by semester, trimester or even months.

The quantities observed in the development period j with the time of accident i are denoted by $X_{i,j}$ and the ones that we are going to study are:

- $X_{i,j} = C_{i,j}$ consists in the payments made in the development period j regarding the claims arisen in time i , in other words the incremental payments;
- $X_{i,j} = A_{i,j}$ consists in the payments made in the development period j , included, regarding the claims arisen in time i , i.e. the cumulative payments;

The factor given by $A_{i,j}$ is easily obtained by

$$A_{i,j} = \sum_{k=0}^j C_{i,k}. \quad (2.1)$$

As we said before, the main goal is to estimate the compensation that the insurer has to pay to the insured. So that, with the given data, it is possible to provide the current liability relative to the claims that occurred in year i , and it will be denoted by R_i , given by

$$R_i = A_{i,\infty} - A_{i,n-i}, \quad 0 \leq i \leq n,$$

where $A_{i,n-i}$ is the total paid amount from year i until the most recent year in consideration of the claims.

Consequently the total amount for the provision will be given by

$$R = \sum_{i=0}^n R_i.$$

2.3 Estimation Methods

Through the nature of claims and the historical data of insurance companies, the estimates for claim reserves can be obtained using a case by case approach or statistical methods.

The case by case approach is used to estimate reserves individually for each claim process and demands to know the claim in hand. This is more suitable for situations where there is a lack of proper data or said data doesn't have appropriate features to use statistical methods.

Whenever the number of claims in an insurer portfolio is big enough, statistical methods are used to estimate the claim reserves. The choice of the method is made

through the assessment of the impact that several external and internal elements might have in the estimation, such as, data homogeneity and statistical nature of estimators.

In order to foresee the total amount of claims that happened in each year and, thus, the provision amount to hold today we need to estimate the lower triangle of the *Run-Off Triangle* in Figure 2.2. In order to do so, we are going to use some usual methods such as the Chain Ladder and Bootstrap methods.

2.4 Confidence Intervals Estimation

When estimating the amount of reserves we can use both deterministic and stochastic models, but with the latter we can also quantify the degree of uncertainty of these estimations and with these error measurements we can assess the quality of the models' adjustment, thus scale the reliability of the attained results. Through this, we can build confidence intervals, in order to see the variance of the reserves estimates, and it will allow us to do as a strict analysis as we want. Usually, we use as an error measurement the Mean Squared Error (*MSE*) which will be defined later.

In order to build a confidence interval for an estimation of the total reserve, $\hat{R}$, we need to know the distribution of the random variable R , which usually is a very difficult, and sometimes, an impossible task. So, in practice, through the Central Limit Theorem, it is assumed that R follows an asymptotically Normal Distribution with mean μ and variance σ^2 ,

$$R \sim N(\mu, \sigma^2),$$

where $\mu = \hat{R}$ and $\sigma^2 = \widehat{MSE}(\hat{R})$. Here, $\hat{R}$ is the estimation of the total reserve and $\widehat{MSE}(\hat{R})$ its mean squared error given by the usage of the stochastic model.

In these conditions, the confidence interval for the estimation of $\hat{R}$ for a confidence level of $1 - \alpha$ is given by:

$$\left[\hat{R} - \Phi_{1-\frac{\alpha}{2}} \cdot \sqrt{\widehat{MSE}(\hat{R})}; \hat{R} + \Phi_{1-\frac{\alpha}{2}} \cdot \sqrt{\widehat{MSE}(\hat{R})} \right] \quad (2.2)$$

where, $\Phi_{1-\frac{\alpha}{2}}$ is the probability quantile of $1 - \frac{\alpha}{2}$ from the Normal Distribution (0,1).

Sometimes, when the inferior limit of the confidence interval is negative or the Standard Error (*SE*) is significantly high, that is

$$\widehat{SE}(\hat{R}) = \sqrt{\widehat{MSE}(\hat{R})} > 0.5 \cdot \hat{R},$$

we have to consider that the random variable R follows a Lognormal Distribution, with μ and σ^2 as parameters, whose first two moments are the same as the estimated values by the model,

$$R \sim \text{Lognormal}(\mu, \sigma^2)$$

where, $E(R) = e^{\mu + \frac{\sigma^2}{2}}$ and $V(R) = e^{2\mu + \sigma^2} \cdot (e^{\sigma^2} - 1)$. This means that the parameters of the distribution are given by

$$\begin{aligned} \begin{cases} E(R) = \hat{R} \\ V(R) = MSE(\hat{R}) \end{cases} &\Leftrightarrow \begin{cases} \mu + \frac{\sigma^2}{2} = \ln(\hat{R}) \\ e^{2\mu} \cdot e^{\sigma^2} \cdot (e^{\sigma^2} - 1) = \widehat{MSE}(\hat{R}) \end{cases} \\ &\Leftrightarrow \begin{cases} \mu = \ln(\hat{R}) - \frac{\sigma^2}{2} \\ e^{2(\ln(\hat{R}) - \frac{\sigma^2}{2})} \cdot e^{\sigma^2} \cdot (e^{\sigma^2} - 1) = \widehat{MSE}(\hat{R}) \end{cases} \\ &\Leftrightarrow \begin{cases} \mu = \ln(\hat{R}) - \frac{\sigma^2}{2} \\ \sigma^2 = \ln\left(\frac{\widehat{MSE}(\hat{R})}{\hat{R}^2} + 1\right) \end{cases} \end{aligned}$$

In this case the confidence interval will be given by

$$\left[e^{\mu - \Phi_{1-\frac{\alpha}{2}} \cdot \sqrt{\sigma^2}}; e^{\mu + \Phi_{1-\frac{\alpha}{2}} \cdot \sqrt{\sigma^2}} \right]$$

We can, in a similar way, define the confidence intervals for the estimations of $\hat{R}_i, 0 \leq i \leq n$, as well for the unknown amounts $C_{i,j}$ or $A_{i,j}, (i, j) : n - i + 1 \leq j \leq n + 1$.

CHAIN LADDER METHOD/ THOMAS MACK

The Chain Ladder is considered to be the most famous method to estimate claim reserves. This is because it is a simple approach and it is distribution-free, meaning that there are little assumptions to consider, in this case there are three assumptions. Complementary to this, it is also known that reserves estimated using the Chain Ladder method for the most recent accident years are very sensitive to variations in the observed data, Mack, 1993.

Regarding the assumptions for the Chain Ladder method, the first one states that there are development factors that represent how claims from a certain year of accident evolve in each development year. We aim to estimate these development factors the best we can for each accident year and the ones regarding the most recent year in our data allows us to establish a total payment amount, after all possible years of development until that recent year.

The second assumption comes due to the fact that this algorithm does not take into account any dependence between accident years, we can additionally assume that the cumulative payments of different accident years are independent. However, in practice, the independence of the accident years can be skewed by certain calendar year effects like major changes in claims handling or in case reserving.

The last, and third, assumption underlying the Chain Ladder method is associated with the variance of the estimators for $C_{i,j}$. Those that should be considered are the ones with the lowest variance.

When studying this method, Thomas Mack developed a stochastic model, with similar results to the traditional and deterministic Chain Ladder method, which has, as an advantage, the means to estimate the mean squared root that explains the variation of results.

Since we are working with claims from insurance of accidents at work we have to take into account a tail factor, i.e. we have to consider that the amount $A_{i,n}$ is not the last claim and the last development factor is higher than 1.

3.1 Model Assumptions

The Chain Ladder method seeks to estimate the expected value of the total cost of claims for each accident year, $A_{i,j}, 0 \leq i \leq n, 0 \leq j \leq n$ and therefore estimate the outstanding claims reserve, $\hat{R}_i$ given by,

$$\hat{R}_i = \hat{A}_{i,\infty} - A_{i,n-i}$$

and subsequently the amount of the total reserve,

$$\hat{R} = \sum_{i=0}^n \hat{R}_i$$

We can obtain the estimated $\hat{A}_{i,\infty}$ by applying the development factors, that will be explained in the next Subsection 3.1.1, to the known payments and estimate the future unknown payments.

3.1.1 First Assumption

As aforementioned, the first Chain Ladder assumption is that there are development factors, $f_j > 0, \forall j \in \{0, \dots, n\}$, that represent how claims from a certain year of accident evolve in each development year, such as

$$E(A_{i,j+1} | A_{i,0}, \dots, A_{i,j}) = A_{i,j} \cdot f_j \quad 0 \leq i \leq n, 0 \leq j \leq n \quad (3.1)$$

where we estimate these development factors by

$$\hat{f}_j = \frac{\sum_{i=0}^{n-j-1} A_{i,j+1}}{\sum_{i=0}^{n-j-1} A_{i,j}}, \quad 0 \leq j \leq n-1 \text{ and } \hat{f}_n = \frac{\hat{A}_{0,\infty}}{A_{0,n}} \quad (3.2)$$

or, equivalently,

$$\hat{f}_j = \sum_{i=0}^{n-j-1} \frac{A_{i,j}}{\sum_{k=0}^{n-j-1} A_{k,j}} \cdot \frac{A_{i,j+1}}{A_{i,j}}, \quad 0 \leq j \leq n-1 \text{ and } \hat{f}_n = \frac{\hat{A}_{0,\infty}}{A_{0,n}}.$$

This first assumption shows the proportionality between the columns of the development triangle showed in equation (3.1) that can also be written by

$$E\left(\frac{A_{i,j+1}}{A_{i,j}} | A_{i,0}, \dots, A_{i,j}\right) = f_j, \quad 0 \leq i \leq n, 0 \leq j \leq n$$

because, when using this condition to compute the expected value, $A_{i,j}$ is a scalar. From this, we can notice that the evolution of the previous known amounts $A_{i,j}$ and the previous development factor $A_{i,j}/A_{i,j-1}$ are independent from the expected value

of the individual development factor $A_{i,j+1}/A_{i,j}$ given by f_j . That is to say that there is no correlation between the individual development factors $A_{i,j}/A_{i,j-1}$ and $A_{i,j+1}/A_{i,j}$, which is surprising since they depend on the same data.

As said before, we will need a tail factor, f_n to correctly represent the future amounts after n development years of a certain accident year i . We can consider various approaches to estimate f_n , such as, bringing the values of $\ln(\hat{f}_j - 1)$ through a line given by $y = a \cdot j + b$, $a < 0$, considering that its estimation will be given by

$$\hat{f}_n = \prod_{j=n}^{\infty} \hat{f}_j$$

Regardless the way this estimation is computed it is essential to analyse the obtained result to ensure that its value is reasonable and represents what is expected in reality.

3.1.2 Second Assumption

The second assumption underlying the Chain Ladder method is a consequence from the first assumption, and it states that

$$\{A_{i,0}, \dots, A_{i,n}\}, \{A_{j,0}, \dots, A_{j,n}\}, i \neq j, \text{ are independent.} \quad (3.3)$$

Through both assumptions we can state that the estimates of the development factors, $\hat{f}_j, 0 \leq j \leq n-1$, are unbiased and uncorrelated, Mack, 1993.

3.1.3 Third Assumption

Now that we have the unbiased estimators we should only consider the ones that have the lowest variance. With this comes the third, and last, assumption underlying the Chain Ladder Method that says given the fact that $\hat{f}_j$ is the $A_{i,j}$ -weighted mean of the individual development factors $A_{i,j+1}/A_{i,j}, 0 \leq i \leq n-j-1$, we can state that $\text{Var}(A_{i,j+1}/A_{i,j} | A_{i,0}, \dots, A_{i,j})$ should be inversely proportional to $A_{i,j}$, or analogously

$$\text{Var}(A_{i,j+1} | A_{i,0}, \dots, A_{i,j}) = A_{i,j} \cdot \sigma_j^2 \quad 0 \leq i \leq n, 0 \leq j \leq n \quad (3.4)$$

where σ_j^2 is the proportionality constant and unknown.

3.2 Mean Squared Error and Standard Error

In order to make sure that the estimation of the total amount of claims given by the Chain Ladder method are reliable, we need to assess the values of the estimators $\hat{A}_{i,n}$ and $\hat{R}_i$ to ensure that they resemble, on average, to their real values $A_{i,n}$ and R_i , because there is a small probability for them to be equal, as well as get to know, also on average, the error of this estimation.

Hence, we can use the Mean Squared Error (*MSE*) to study this variability of the estimations regarding its future randomness since what happened in the past does not concern us.

The MSE of the estimator $\hat{A}_{i,n}$ is defined by Mack, Mack, 1993, as

$$\begin{aligned} MSE(\hat{R}_i) &= MSE(\hat{A}_{i,n}) = E\left(\left(\hat{A}_{i,n} - A_{i,n}\right)^2 | D\right) \\ &= Var(A_{i,n} | D) + \left(E(A_{i,n} | D) - \hat{A}_{i,n}\right)^2 \end{aligned}$$

where $D = \{A_{i,j} : i + j \leq n + 1\}$ is the set all observed data so far.

From the underlying assumptions (3.1), (3.3) and (3.4), the Mean Squared Error of $\hat{R}_i$ can be estimated as,

$$\widehat{MSE}(\hat{R}_i) = \hat{A}_{i,n}^2 \sum_{j=n+1-i}^{n-1} \frac{\hat{\sigma}_j^2}{\hat{f}_j^2} \left(\frac{1}{\hat{A}_{i,j}} + \frac{1}{\sum_{k=1}^{n-j} A_{k,j}} \right) \quad (3.5)$$

Mack, Mack, 1993, defines the estimator of σ_j^2 as,

$$\hat{\sigma}_j^2 = \frac{1}{n-j-1} \sum_{i=1}^{n-j} A_{i,j} \left(\frac{A_{i,j+1}}{A_{i,j}} - \hat{f}_j \right)^2, \quad 1 \leq j \leq n-2 \quad (3.6)$$

as an unbiased estimator of σ_j^2 , $1 \leq j \leq n-2$. The estimation of σ_{n-1}^2 is still missing and Mack, Mack, 1993, suggests to add an additional member by extrapolating the usually exponentially decreasing series $\hat{\sigma}_1, \dots, \hat{\sigma}_{n-3}, \hat{\sigma}_{n-2}$ by loglinear regression, for example, by requiring that

$$\frac{\hat{\sigma}_{n-3}}{\hat{\sigma}_{n-2}} = \frac{\hat{\sigma}_{n-2}}{\hat{\sigma}_{n-1}}$$

to be true as long as $\hat{\sigma}_{n-3} > \hat{\sigma}_{n-2}$, which leads to,

$$\hat{\sigma}_{n-1}^2 = \min\left(\frac{\hat{\sigma}_{n-2}^4}{\hat{\sigma}_{n-3}^2}, \min(\hat{\sigma}_{n-3}^2, \hat{\sigma}_{n-2}^2)\right) \quad (3.7)$$

3.3 Verification of Assumptions

In Section 3.1 we presented the assumptions underlying the Thomas Mack Model, and it is important that our data follows these assumptions. To make sure this happens we will do some statistical tests for each assumption to our data.

When we can prove that our data follows all the assumptions underlying the Thomas Mack Model we can start applying the Model to our data.

3.3.1 First Assumption

The first assumption is described by (3.1), and for a fixed j , it conveys a linear relation between $A_{i,j+1}$ and $A_{i,j}$ with a slope given by f_j which can be perceived as a linear regression model given by

$$y_i = mx_i + b + \varepsilon_i$$

where the regression's coefficients are $b = 0$ and $m = f_j$ and its residuals are ε_i with $E(\varepsilon_i) = 0$, so,

$$E(y_i) = mx_i + b$$

To see this linear relation we must analyse some plots from our data, where we draw a line that crosses the origin and has a slope of f_j , so the ordered pair $(A_{i,j}, A_{i,j+1})$ should draw near that line. In case there are significant deviations the necessity of searching for new estimators of f_j arises, in order to have a better adjustment to our data, if this is not possible, then we should reject the application of the model to this data.

The first assumption, as mentioned before, also requires for the individual development factors, $A_{i,j}/A_{i,j-1}$ and $A_{i,j+1}/A_{i,j}$, to be uncorrelated. This can be verified using the Spearman Test which, even though it is just an approximation, it does not use any distribution for the data.

To apply this test we consider a fixed development year $j, 1 \leq j \leq n-1$, and sort the individual factors $A_{i,j+1}/A_{i,j}$ in ascending order, denoting by $r_{i,j}$ the rank given to $A_{i,j+1}/A_{i,j}, 1 \leq r_{i,j} \leq n-j$. Similarly, we sort the precedent factors $A_{i,j}/A_{i,j-1}$, dismissing the last factor $A_{n-j,j}/A_{n-j,j-1}$, in ascending order, denoting by $s_{i,j}$ the rank given by this type of factor, $1 \leq s_{i,j} \leq n-j$. Putting this into a hypothesis test we have the following hypothesis:

H_0 : The development factors are not correlated

vs

H_1 : The development factors are correlated

The Spearman's correlation coefficient, T_j , is given by

$$T_j = 1 - 6 \cdot \sum_{i=0}^{n-j-1} \frac{(r_{i,j} - s_{i,j})^2}{(n-j)^3 - n + j}, \quad 1 \leq j \leq n-2 \quad (3.8)$$

where $-1 \leq T_j \leq 1$.

In the absence of correlation, i.e., under the null hypothesis we have

$$E(T_j) = 0 \text{ and } V(T_j) = \frac{1}{n-j-1}$$

So if T_j takes a value close to zero it means that the development factors used to sort $r_{i,j}$ and $s_{i,j}$ are not correlated, and every other value of T_j suggests a positive or negative correlation.

In order to compute the Spearman Test we need to consider the development triangle as a whole, to take into account all the correlation that might exist, and for this we need to compute the weighted average of the T_j 's, where, in order to have a estimator of minimum variance, the weights are inversely proportional to $V(T_j)$. So,

$$\begin{aligned} T &= \frac{\sum_{j=1}^{n-2} (n-j-1) \cdot T_j}{\sum_{j=1}^{n-2} n-j-1} = \\ &= \sum_{j=1}^{n-2} \frac{n-j-1}{\frac{(n-1) \cdot (n-2)}{2}} \cdot T_j, \end{aligned}$$

where in the absence of correlation, i.e., under the null hypothesis we have

$$E(T) = 0 \text{ and } V(T) = \frac{1}{\frac{(n-1) \cdot (n-2)}{2}}$$

Thereupon, it can be assumed that T follows a Normal distribution. This is due to the approximation of T_j 's to this distribution with $n-j \geq 10$ and that T results from the weighted average of the several, non-correlated, T_j 's.

Consequently, we want to reject our alternative hypothesis of correlation within the development factors. So our estimation of T must belong of the confidence interval of 50% (this is the significance level used, instead of 95%, because this test is only an approximation and we only want to find correlations in a relevant part of the *Run-off Triangle*) in order to reject this hypothesis. The confidence interval is given by

$$\left[-\Phi_{1-\frac{\alpha}{2}}^{-1} \sqrt{V(T)}; \Phi_{1-\frac{\alpha}{2}}^{-1} \sqrt{V(T)} \right]$$

If and when T does not belong to this confidence interval the correlations must be analysed with greater detail and, possibly, choose to not apply this method.

3.3.2 Second Assumption

The second assumption, as we saw above in (3.3), states that the development factors are independent for all of the development years.

In order to verify this assumption the following test is presented.

Considering that any occurrence could influence one of these diagonals,

$$D_k = \{A_{k,0}, A_{k-1,1}, \dots, A_{0,k}\} \quad 0 \leq k \leq n$$

hence, equally influence the adjacent individual development factors

$$A_k = \left\{ \frac{A_{k,1}}{A_{k,0}}, \frac{A_{k-1,2}}{A_{k-1,1}}, \dots, \frac{A_{0,k+1}}{A_{0,k}} \right\} \text{ and } A_{k-1} = \left\{ \frac{A_{k-1,1}}{A_{k-1,0}}, \frac{A_{k-2,2}}{A_{k-2,1}}, \dots, \frac{A_{0,k}}{A_{0,k-1}} \right\}$$

We must, then, split all of the individual development factors in two sets so we can analyse the independence between the occurrence years. These two sets will be composed by the highest and the lowest development factors and afterwards we have to observe the matrix diagonal, to see if there is a predominance of the elements in one of these sets.

With this in mind we start by sorting the elements for each set

$$F_j = \left\{ \frac{A_{i,j+1}}{A_{i,j}} \mid 0 \leq i \leq n-1-j \right\} \quad 0 \leq j \leq n-1$$

For a fixed j , F_j has all the observed individual development factors from j to $j+1$ years.

For each set of F_j , we set apart the elements in two subsets where LF_j has the half largest individual development factors and SF_j the other half of the smallest individual development factors. Note that the number of elements of each subset must be equal, and in case the set F_j has an odd amount of elements then we must ignore the median.

From this, we know that each individual development factor will either be in one of the subsets or overlooked.

$$L = LF_0 \cup LF_1 \cup \dots \cup LF_{n-2}$$

$$S = SF_0 \cup SF_1 \cup \dots \cup SF_{n-2}$$

Thus, for a given individual development factor the probability of belonging to the subset S or L is 0.5.

To find how many elements belong to the subset L (large factors) and to the subset S (small factors) we must count the elements for each diagonal A_k of the development factors. We expect to find nearly the same amount in both subsets if there are no effects of a given year.

Considering L_k and S_k the number of elements of L and S that belong to A_k , respectively. If there is a prominence of large or small factors in the k diagonal then $Z_k = \min(L_k, S_k)$ is substantially smaller than $(L_k + S_k)/2$ meaning that the second assumption (3.3) is rejected.

Taking into account the verification of (3.3) we know that L_k and S_k follow a Binomial distribution where $a = L_k + S_k$ and $p = 0.5$. So, assuming an approximation to the Normal distribution we reject the following null hypothesis

H_0 : The occurrence years are independent

vs

H_1 : The occurrence years are not independent

with a significance level of 5%, every time the following does not happen

$$Z \in \left[E(Z) - \Phi_{1-\frac{\alpha}{2}}^{-1} \cdot \sqrt{V(Z)}; E(Z) + \Phi_{1-\frac{\alpha}{2}}^{-1} \cdot \sqrt{V(Z)} \right] \quad (3.9)$$

where Z is the random variable given by $Z = Z_1 + Z_2 + \dots + Z_{k-1}$, and under the null hypothesis we have that

$$E(Z) = \sum_k E(Z_k) \text{ and } V(Z) = \sum_k V(Z_k)$$

with

$$E(Z_k) = \frac{a}{2} - \binom{a-1}{v} \frac{a}{2^a}$$

and

$$V(Z_k) = \frac{a(a-1)}{4} - \binom{a-1}{v} + E(Z_k) - (E(Z_k))^2,$$

where v is the biggest integer less or equal to $(a-1)/2$.

Summarizing the aforementioned, if Z does not fall in the confidence interval above we cannot proceed with the implementation of the model.

3.3.3 Third Assumption

To verify the third assumption given by (3.4) we need to take into consideration that the conditional variance of $A_{i,j+1}$ is directly proportional to $A_{i,j}$ with σ_j^2 as a constant of proportionality. With the known section of the *Run-off Triangle* we can estimate $V(A_{i,j+1}|A_{i,0}, \dots, A_{i,j})$ through

$$\begin{aligned} E \left[\left(A_{i,j+1} - E(A_{i,j+1}|A_{i,0}, \dots, A_{i,j}) \right)^2 \right] &\stackrel{(3.1)}{=} \\ &\stackrel{(3.1)}{=} \left(A_{i,j+1} - A_{i,j} \cdot f_j \right)^2 \approx \\ &\approx \left(A_{i,j+1} - A_{i,j} \cdot \hat{f}_j \right)^2, \\ &0 \leq i \leq n-1, 0 \leq j \leq n-1 \end{aligned}$$

The estimated values for the conditional variance should come closer to $A_{i,j} \cdot \sigma_j^2$, if (3.4) is verified, i.e., the obtained residuals will also be proportional, or come close, to $A_{i,j}$.

If (3.4) verifies we will have

$$\left(A_{i,j+1} - A_{i,j} \cdot \hat{f}_j\right)^2 \approx A_{i,j} \cdot \sigma_j^2$$

or

$$\frac{A_{i,j+1} - A_{i,j} \cdot \hat{f}_k}{\sqrt{A_{i,j}}} \approx \sigma_j$$

When obtained through this, the values will no longer be linked to $A_{i,j}$, so it will be enough to represent the ordered pair

$$\left(\frac{A_{i,j+1} - A_{i,j} \cdot \hat{f}_j}{\sqrt{A_{i,j}}}, A_{i,j} \right), \text{ for a fixed } j \quad (3.10)$$

in a plot to see if the set of obtained points do not present any type of trend. If this happens, then we can accept the assumption given by (3.4) as valid.

However, if there is a trend the assumption will not be verified and consequently we must get other f_j estimators, or, if necessary reject the implementation of this model.

3.4 Practical Application

Let's consider the following *Run-off Triangle* as our development triangle with incremental paid claims from accidents at work, since 2007 until 2016. All values of our data are in thousands of Euros and all the results that we will see below were obtained using the version 0.2.14 of the ChainLadder package of RCrn, Gesmann et al., 2021.

[illegible]

Figure 3.1: Incremental paid claims

To apply the Thomas Mack Model to our data we need to change the *Run-Off Triangle* in order to have cumulative paid claims.

	dev									
origin	0	1	2	3	4	5	6	7	8	9
2007	182487.5	283428.5	298249.6	303722.8	306458.5	306789.3	307263.7	308719.3	309515.5	310253.8
2008	202627.4	306135.3	322697.8	327967.6	330973.7	332978.8	334047.3	334591.8	335716.8	
2009	188765.5	288300.6	303118.8	308136.5	311664.4	313679.4	315021.5	315916.6		
2010	190072.7	285882.6	301218.2	307229.7	310142.9	311193.2	312095.4			
2011	188679.2	282818.6	297873.5	303266.2	306004.5	307831.0				
2012	178899.8	265004.8	278476.9	284306.6	287374.1					
2013	178266.6	264024.5	279587.7	286366.3						
2014	184083.4	275485.0	291413.2							
2015	192314.3	287536.1								
2016	193896.1									

Figure 3.2: Cumulative paid claims

The first thing we need to do is to verify if our data follows all the assumptions implied in the Model and for that we need to calculate the development factors, f_j , given by (3.2) and are the following

	1	2	3	4	5	6	7	8	9
Development factors	1.505528	1.053999	1.01911	1.009805	1.004618	1.002995	1.003027	1.002986	1.002385

Figure 3.3: Estimation of Development Factors

Using the linear regression explained at the end of 3.1.1, we've obtained a tail factor of 1.001286 and since this is close to 1 we will not be using it, therefore the *Run-off Triangle* in Figure 3.1 will be considered as closed.

As said in Subsection 3.3.1, we will start by analysing the following plots to see the linear relation underlying the assumption given by (3.1).

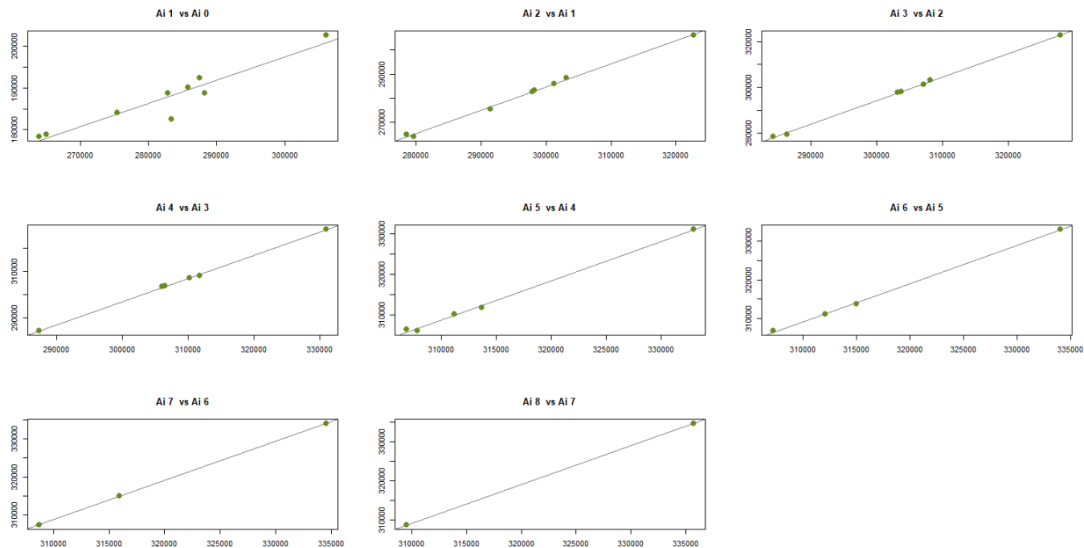


Figure 3.4: Data adjustment to f_j

As we can see in Figure 3.4, for each column j , the line with slope of f_j seems to be well adjusted to the data and for that, we can assume that our data follows the assumption in (3.1), which says that there is proportionality between adjacent development years, .

The other characteristic of the first assumption that we need to prove is that the individual development factors $A_{i,j}/A_{i,j-1}$ and $A_{i,j+1}/A_{i,j}$ are not correlated by applying a Spearman test to these individual development factors. By computing $A_{i,j+1}/A_{i,j}$ for each entry of the *Run-off Triangle* and we get the following output.

origin	dev	1	2	3	4	5	6	7	8	9	10
1	1.553139	1.052292	1.018351	1.009007	1.001079	1.001546	1.004737	1.002579	1.002385		
2	1.510829	1.054102	1.016330	1.009166	1.006058	1.003209	1.001630	1.003362			
3	1.527295	1.051398	1.016553	1.011449	1.006465	1.004279	1.002841				
4	1.504070	1.053643	1.019958	1.009482	1.003386	1.002899					
5	1.498939	1.053232	1.018104	1.009029	1.005969						
6	1.481303	1.050837	1.020934	1.010789							
7	1.481065	1.058946	1.024245								
8	1.496523	1.057819									
9	1.495137										
10											

Figure 3.5: Individual development factors

To do the Spearman test we must, as explained above, rank the individual development factors in two different ways in order to build $r_{i,j}$ and $s_{i,j}$.

origin	dev	1	2	3	4	5	6	7	8	9	10
1	3	4	1	1	1	3	1	1			
2	6	1	3	4	3	1	2				
3	2	2	6	5	4	2					
4	5	5	4	2	2						
5	4	3	2	3							
6	1	6	5								
7	8	7									
8	7										
9											
10											

(a) $r_{i,j}$

origin	dev	1	2	3	4	5	6	7	8	9	10
1	8	3	4	1	1	1	2	1			
2	6	6	1	3	3	2	1				
3	7	2	2	5	4	3					
4	5	5	5	4	2						
5	4	4	3	2							
6	2	1	6								
7	1	7									
8	3										
9											
10											

(b) $s_{i,j}$

Figure 3.6: Individual development factors ranked

Using (3.5) we get, that for each j , the estimations of the Spearman's correlation coefficient is

Tk's	1	2	3	4	5	6	7	8	9
	-0.38095238	0.07142857	0.08571429	0.70000000	1.00000000	-0.50000000	-1.00000000		

Figure 3.7: Individual development factors

And the final estimation for T is 0.07108844.

So, as we said before, when there is no correlation it is expected for

$$E(T) = 0 \text{ and } V(T) = \frac{1}{28} = 0.03571429$$

therefore, the confidence interval of 50% for the estimation of T will be

$$[-0.1274666; 0.1274666]$$

Since our estimative belongs to this interval we will not reject the null hypothesis with a level of statistical significance, meaning that we do not reject the hypothesis of the non-correlation for the development factors.

Next, we will prove if our development factors are independent for all of the development years as stated by (3.3), and for this we will implement what was described in Subsection 3.3.2.

We already have the *Run-Off Triangle* with the individual development factors, Figure 3.5, so we only need to rank each individual development factor by column from the smallest to the largest.

	dev									
origin	1	2	3	4	5	6	7	8	9	10
1	9	3	4	1	1	1	3	1	1	
2	7	6	1	3	4	3	1	2		
3	8	2	2	6	5	4	2			
4	6	5	5	4	2	2				
5	5	4	3	2	3					
6	2	1	6	5						
7	1	8	7							
8	4	7								
9	3									
10										

Figure 3.8: Individual development factors ranked from smallest to largest

As aforementioned, the ranked individual development factors are assigned to the L or S sets, if they are above or below the median, respectively, and are labelled with a " L " or " S ". When an element is equal to the median it is labelled with an "*" and is ignored, meaning it does belong to neither of the sets.

	dev									
origin	1	2	3	4	5	6	7	8	9	10
1	"L"	"S"	"*"	"S"	"S"	"S"	"L"	"S"	"*"	
2	"L"	"L"	"S"	"S"	"L"	"L"	"S"	"L"		
3	"L"	"S"	"S"	"L"	"L"	"L"	"*"			
4	"L"	"L"	"L"	"L"	"S"	"S"				
5	"*"	"S"	"S"	"S"	"S"	"*"				
6	"S"	"S"	"L"	"L"						
7	"S"	"L"	"L"							
8	"S"	"L"								
9	"S"									
10										

Figure 3.9: Elements of the L and S sets

Let S_j and L_j be the number of elements belonging to the S and L sets, respectively, from a j diagonal.

j	S_j	L_j	Z_j	v	a	$E(Z_j)$	$E(Z_j)$
0	0	1	0	1	0	0	0
1	1	1	1	2	0	0.5	0.25
2	0	2	0	2	0	0.5	0.25
3	3	1	1	4	1	1.25	0.44
4	3	1	1	4	1	1.25	0.44
5	3	3	3	6	2	2.06	0.62
6	3	4	3	7	3	2.41	0.55
7	5	3	3	8	3	2.91	0.80
8	2	4	2	6	2	2.06	0.62
Total			14			12.94	3.97

Table 3.1: Results of Z_j and Z variables

From the results on Table 3.1 we can calculate the lower and higher limit of our confidence interval given by (3.9),

$$[10.80648, 18.60758]$$

Since, our Z is in the confidence interval we don't reject the null hypothesis, meaning we don't have any statistical reason to state that the occurrence years are not independent.

For the third assumption, we will look into a series of plots that represent the ordered pair in (3.10) for each j .

Thus, for our dataset the plots for the points given by (3.10) are the following,

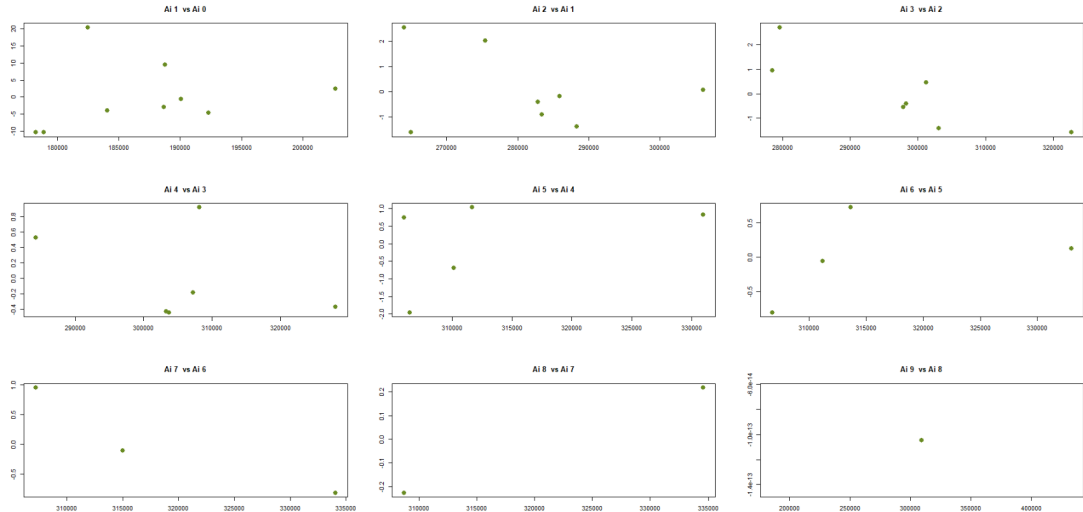


Figure 3.10: Weighted Residuals

As we can see, it appears that there is no type of trend, meaning that the assumption (3.4) is valid.

Now that all the assumptions underlying the Thomas Mack Model have been proved we can, therefore, begin to apply the Model to our data.

As aforementioned, the Model estimates the future amounts $A_{i,j}$, thus the provisions estimates R_i . The estimations of the development factors, $\hat{f}_j$, calculated resorting to (3.2) are represented in Figure 3.3.

Thus, the lower triangle can now be completed using (3.1), which gives us,

origin	dev	1	2	3	4	5	6	7	8	9	10
1	184770.3	95690.27	15200.62	5952.762	3433.424	1378.489	831.3374	378.6499	857.0619	534.9824	
2	204373.1	100804.36	16487.38	6084.441	3147.758	1517.709	1268.0730	1524.8842	897.4392		
3	190889.0	95147.77	15703.64	5395.558	2036.987	1424.871	1731.8075	1104.8151			
4	191514.9	99307.43	17099.08	6072.925	3239.707	1147.308	866.9459				
5	182260.2	93355.76	13016.83	5653.455	2218.485	2213.445					
6	173797.8	89338.81	13919.94	5389.402	3425.463						
7	175869.1	85368.27	14396.12	4738.135							
8	188362.6	93995.82	14464.89								
9	191485.5	94880.99									
10	204358.0										

Figure 3.11: Full Triangle with estimated future claim amounts

Thereafter, the estimated reserve for each accident year is

Reserve of each i	1	2	3	4	5	6	7	8	9	10
	0	800.7296	1699.229	2628.602	3522.318	4630.394	7467.24	13313.1	29371.94	127839.3

Figure 3.12: Estimated reserve by Accident Year

Consequently the total reserve is 195285.8.

Now to measure the variability of $\hat{R}_i$ using (3.5), we need to estimate σ_j^2 using (3.6), so

Sigma squared for each j	1	2	3	4	5	6	7	8
	92.99855	1.727421	1.015336	0.2749276	1.514359	0.3920078	0.7753262	0.05117787

Figure 3.13: Estimation of σ_j^2 , $1 \leq j \leq 8$

Since $0.7753622 > 0.05117787$, we are in conditions to use (3.7) to extrapolate the last estimator, $\hat{\sigma}_{n-1}$, which value is 0.003378157.

Finally, the variability of the estimators for R_i given by (3.5) is

MSE for each Ri	1	2	3	4	5	6	7	8	9
	3906945026	2035627	655795.1	47313.22	1415514	90927.69	361559.8	1675.252	9.539173

Figure 3.14: Estimation of $\widehat{MSE}(\hat{R}_i)$

From Section 2.4 and from the estimations above, we get the estimation of the confidence intervals given by (2.2) for each occurrence year.

	2007	2008	2009	2010	2011	2012	2013	2014	2015	2016
Lower Limit	0	618.1400	1241.568	1419.543	2097.650	2609.009	5334.218	10526.44	26015.77	117451.8
Upper Limit	0	983.3192	2156.890	3837.662	4946.986	6651.779	9600.261	16099.77	32728.12	138226.8

Figure 3.15: Estimation of confidence interval for R_i

BOOTSTRAP METHOD

The Bootstrap method was introduced by Brandlen Efron, in 1979, Taylor, 2012, and its essence is to take a sample of a triangle of paid claims from an unknown distribution in order to obtain information about future claim payments, or reserve, by creating large amount of sets of pseudo-data through re-sampling with replacement from the observed data in the triangle, England and Verrall, 2002. These sets have the same underlying distribution so we can attain, for each set, statistics of interest and the distribution of those statistics, such as, the mean of each set of pseudo-data, the standard deviation of the set of means and estimate the standard error of the mean.

Initially, Bootstrapping was only used to get estimation errors of reserve estimates from the Chain Ladder Model. Later, some adjustments were made to the Model in order to simulate the process error and allowing to get a predictive distribution.

When we apply the bootstrap, it is expected for the data to be independent and identically distributed, in order for the re-sampling with replacement to happen from the original data. Since we have a regression type problem, our data is independent (as proved in the subsection 3.3.2) but it is not identically distributed. To surpass this problem it is a standard procedure to bootstrap the residuals from the incremental amounts, $C_{i,j}$ instead, because these are approximately independent and identically distributed, with a proper definition for the residuals dismissing the structure of the *Run-Off Triangle*.

Thereafter, we illustrate the main steps of the Bootstrapping

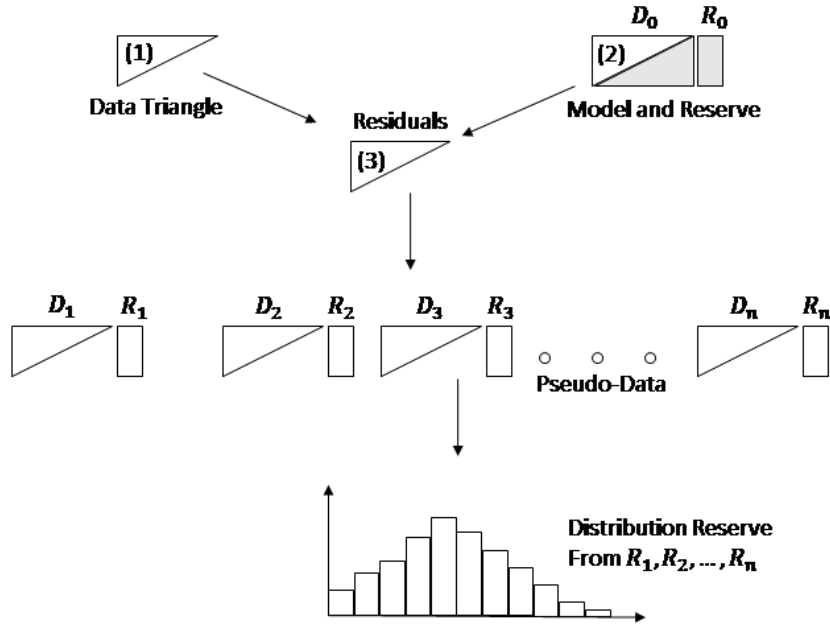


Figure 4.1: Bootstrap Method

The *Data Triangle* is our data and the grey area in *Model and Reserve* is the projection of future payments and the estimated reserves using the Bootstrap method and the triangle denoted by (2) is a fitted model for the past data. The difference between triangles (1) and (2) gives one set of *Residuals*, triangle (3), meaning that our data and the fitted data may differ, Lowe, 1994.

If the data is a implementation of a random process, another execution of the process might lead to another set of data that differs from the model, hence another set of residuals.

When producing various sets of data with residuals re-sampling and adding them to the fitted model we create many sets of possible data triangles, denoted as *Pseudo-Data*. For each triangle the reserving method is implemented, producing a series of reserve estimated *Pseudo-Reserves*. After several iterations we get a large collection of Pseudo-Reserves, which will have a specific distribution with its mean and variance. The latter gives us a measure of variability of the reserve estimate, Lowe, 1994.

4.1 Bootstrap Method

According to, England and Verrall, 2002, we need to prepare our data before bootstrapping. Since the data (Data Triangle (1) in Figure 4.1) does not follow any distribution that we know of, we will be using the Chain Ladder/ Thomas Mack (Chapter 3) to get the development factors from the cumulative payments.

Then, it is required to build a new *Run-off Triangle* (Model and Reserve (2) in Figure 4.1) with cumulative fitted values of the observed data, using a backwards recursion starting from the latest observed cumulative payment in the main diagonal divided by the calculated development factors.

Meaning, the new fitted values, $D_{i,j}$, $0 \leq i \leq n$, $0 \leq j \leq n - i$, are given by

$$D_{i,n-i} = A_{i,n-i} \quad \text{and} \quad D_{i,j-1} = \frac{D_{i,j}}{\hat{f}_j} \quad (4.1)$$

Next, in order to assess the quality of the fitted values we will calculate the unscaled residuals that will be used during the re-sampling using the Pearson residuals $r_{i,j}^P$ given by

$$r_{i,j}^P = \frac{C_{i,j} - \hat{m}_{i,j}}{\sqrt{\hat{m}_{i,j}}} \quad (4.2)$$

where, $C_{i,j}$ are the incremental payments from the observed data and $\hat{m}_{i,j}$ are the incremental fitted payments (4.1).

We now need to adjust the Person residuals (Residuals (3) in 4.1), to replicate the bias correction using an analytic approach, considering the degrees of freedom, i.e. the number of observations minus the number of parameters,

$$r_{i,j}^{Padj} = \sqrt{\frac{n}{\frac{1}{2}n(n+1) - 2n + 1}} \cdot r_{i,j}^P \quad (4.3)$$

Also using the degrees of freedom we can calculate the Pearson scale parameter ϕ ,

$$\phi = \frac{\sum_i^n \sum_j^{n-i} (r_{i,j}^P)^2}{\frac{1}{2}n(n+1) - 2n + 1} \quad (4.4)$$

Now, we will begin the algorithm portion of the Bootstrap Method with an iterative loop where we create a new past *Run-Off Triangle* by re-sampling the adjusted Person residuals with replacement, thus obtaining a set of pseudo-data.

For this first step we need to compute (4.2) for $C_{i,j}$, i.e.,

$$C_{i,j} = r_{i,j}^{Padj} \cdot \sqrt{\hat{m}_{i,j}} + \hat{m}_{i,j}$$

Consecutively, we produce the cumulative form of the pseudo-data and use the Chain Ladder method upon it, generating the estimate of the future cumulative payments.

Then, we go back to the beginning of the iterative loop and create a new pseudo-data with re-sampled adjusted Person residuals with replacement.

For each completed pseudo-data we compute the respective reserve for each occurrence year and thus a total reserve, as we did in the in the Thomas Mack method and store all the information, this will form a predictive distribution. We can also calculate variability measures from the stored information.

4.2 Pratical Application

The *Run-Off Triangle* that we will be analysing with the Bootstrap method represents the incremental claims from 2007 to 2016 of accidents at work insurance in thousand of Euros. This analysis was performed using the 0.2.14 ChainLadder package of the RCrn, Gesmann et al., 2021, and is the same as in 3.4.

	dev									
origin	0	1	2	3	4	5	6	7	8	9
2007	182487.5	100941.02	14821.08	5473.170	2735.748	330.7869	474.3595	1455.6075	796.284	738.2361
2008	202627.4	103507.87	16562.55	5269.781	3006.117	2005.0792	1068.5394	544.5107	1124.944	
2009	188765.5	99535.08	14818.20	5017.671	3527.896	2014.9967	1342.0986	895.1047		
2010	190072.7	95809.91	15335.59	6011.594	2913.137	1050.2703	902.2470			
2011	188679.2	94139.42	15054.90	5392.696	2738.328	1826.4831				
2012	178899.8	86104.97	13472.14	5829.670	3067.442					
2013	178266.6	85757.91	15563.21	6778.554						
2014	184083.4	91401.69	15928.18							
2015	192314.3	95221.83								
2016	193896.1									

Figure 4.2: Incremental paid claims

We start by applying the Chain Ladder method to the following cumulative payments, so we can obtain the development factors of each development year.

	dev									
origin	0	1	2	3	4	5	6	7	8	9
2007	182487.5	283428.5	298249.6	303722.8	306458.5	306789.3	307263.7	308719.3	309515.5	310253.8
2008	202627.4	306135.3	322697.8	327967.6	330973.7	332978.8	334047.3	334591.8	335716.8	
2009	188765.5	288300.6	303118.8	308136.5	311664.4	313679.4	315021.5	315916.6		
2010	190072.7	285882.6	301218.2	307229.7	310142.9	311193.2	312095.4			
2011	188679.2	282818.6	297873.5	303266.2	306004.5	307831.0				
2012	178899.8	265004.8	278476.9	284306.6	287374.1					
2013	178266.6	264024.5	279587.7	286366.3						
2014	184083.4	275485.0	291413.2							
2015	192314.3	287536.1								
2016	193896.1									

Figure 4.3: Cumulative paid claims

Using the proved assumption for this *Run-off Triangle* given by (3.2), we get the development factors below.

	1	2	3	4	5	6	7	8	9
Development factors	1.505528	1.053999	1.01911	1.009805	1.004618	1.002995	1.003027	1.002986	1.002385

Figure 4.4: Estimation of Development Factors

As explained in Section 3.4 we won't be using the tail factor since the estimation was close to 1.

Now that we have the development factors for each development year estimated we can use a backwards recursion to build our fitted values, $\hat{m}_{i,j}$. Starting with our main diagonal, i.e., the last known payments from each occurrence year, we divide them by their respective development factor using (4.1).

origin	dev	1	2	3	4	5	6	7	8	9	10
1		186976.6	281498.6	296699.2	302369.3	305334.0	306743.9	307662.5	308593.9	309515.5	310253.8
2		202804.6	305328.1	321815.5	327965.5	331181.3	332710.5	333706.9	334717.2	335716.8	
3		191413.4	288178.3	303739.6	309544.2	312579.3	314022.7	314963.1	315916.6		
4		189670.6	285554.5	300974.1	306725.9	309733.3	311163.6	312095.4			
5		187639.3	282496.2	297750.7	303440.8	306416.1	307831.0				
6		175978.5	264940.6	279247.1	284583.7	287374.1					
7		177080.9	266600.2	280996.3	286366.3						
8		183645.5	276483.4	291413.2							
9		190986.9	287536.1								
10		193896.1									

Figure 4.5: Fitted Payments

The next step is to calculate the unadjusted Person residuals using (4.2). For the simplicity of the implementation we will be using the following unadjusted residuals shown in the Figure 4.6.

origin	dev	1	2	3	4	5	6	7	8	9	10
1		-1.283216e+01	25.8069160	-3.8049480	-3.231923	-5.198506	-35.52244	-18.117321	21.229469	-5.102527	0
2		-4.865257e-01	3.8000998	0.7235755	-13.874059	-4.568750	15.03976	2.825810	-18.112328	4.899367	
3		-7.480675e+00	11.0073964	-7.3630899	-12.766810	11.056263	18.59804	16.190530	-2.338505		
4		1.140971e+00	-0.2951403	-0.8365404	4.234801	-2.126103	-12.41792	-1.198559			
5		2.967289e+00	-2.8795351	-1.9974242	-4.874103	-5.368931	13.52475				
6		8.607466e+00	-11.8401798	-8.6222053	8.343688	6.483405					
7		3.482921e+00	-15.5391010	12.0229830	23.758906						
8		1.262993e+00	-5.8264186	10.0994633							
9		3.754344e+00	-5.2803341								
10		-8.169504e-14									

Figure 4.6: Unadjusted Person Residuals

Ahead of the implementation of the method, we need to see if our residuals follow the independent assumption through the observation of the following figures.

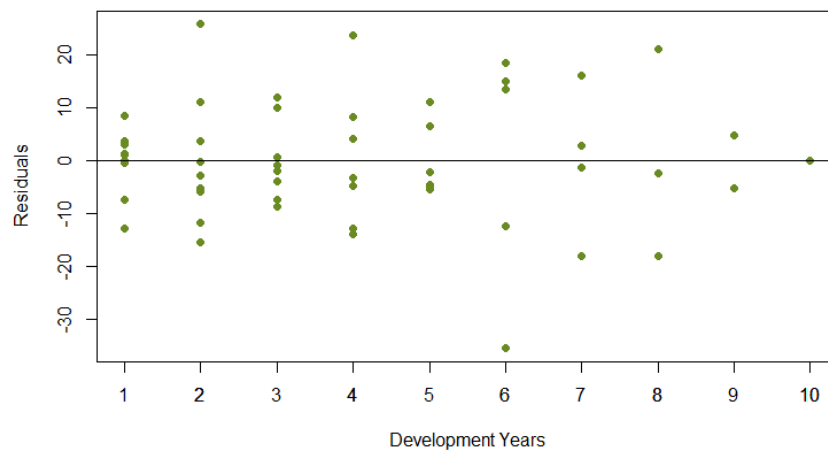


Figure 4.7: Residuals of each Development Year

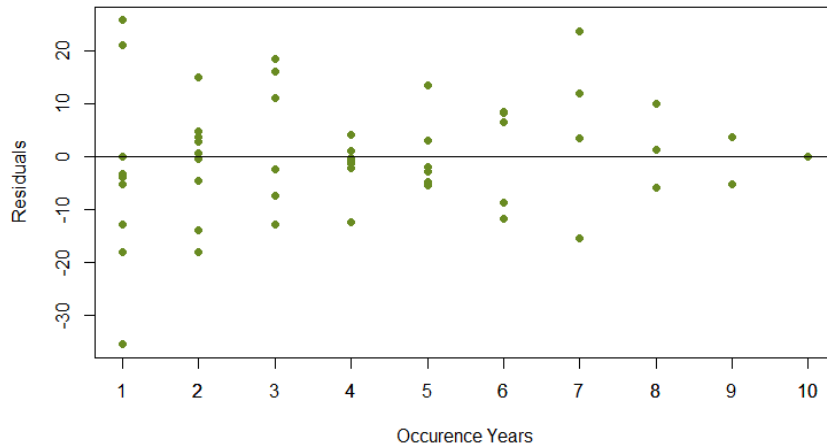


Figure 4.8: Residuals of each Occurrence Year

From the Figures 4.7 and 4.8 we do not see any specific pattern so we can state that these residuals are independent.

Using a process of re-sampling with replacement we can create various sets of simulated residuals, thus creating various sets of pseudo-data. The Figure 4.9 shows an example of a set of simulated residuals.

origin	dev	1	2	3	4	5	6	7	8	9	10
1	13.52475	-5.36893	11.00740	-15.53910	-1.99742	1.26299	-11.84018	-12.83216	4.89937	-0.83654	
2	0.00000	-7.48068	21.22947	-12.41792	-7.48068	23.75891	23.75891	18.59804	0.00000		
3	-3.23192	-35.52244	1.26299	-5.36893	3.75434	8.34369	-15.53910	1.14097			
4	16.19053	2.82581	-7.36309	-1.19856	-5.28033	2.82581	-0.48653				
5	-11.84018	13.52475	11.00740	-12.83216	-0.48653	8.34369					
6	2.96729	18.59804	1.14097	-4.87410	11.05626						
7	-11.84018	-35.52244	-12.83216	1.26299							
8	0.72358	-3.23192	13.52475								
9	-4.87410	-2.33850									
10	-18.11732										

Figure 4.9: Example of simulated residuals

Using the equation (4.2) we can obtain $C_{i,j}$ from the residuals from Figure 4.9, thus creating a set of pseudo-data with incremental payments.

origin	dev	1	2	3	4	5	6	7	8	9	10
1	192824.8	92871.30	16557.73	4499.969	2855.993	1457.323	559.7527	539.7967	1070.3410	715.5069	
2	202804.6	100128.20	19213.32	5176.204	2791.515	2458.358	1746.3343	1601.4063	999.6223		
3	189999.4	85714.88	15718.86	5395.558	3241.936	1760.345	463.8862	988.7587			
4	196721.8	96758.86	14505.31	5660.857	2717.893	1537.081	916.9959				
5	182510.4	99022.40	16614.00	4722.185	2948.720	1728.745					
6	177223.3	94509.25	14442.98	4980.482	3374.398						
7	172098.4	78891.12	12856.47	5462.524							
8	183955.5	91853.21	16582.36								
9	188856.8	95822.62									
10	185918.4										

Figure 4.10: Example of a set of pseudo-data

To this example of a set of pseudo-data, we apply the Chain Ladder method and end up with the following reserves for each occurrence year, which gives us a total reserve amount for this example of pseudo-data of 185450.4.

Reserve of each i 0 769.6034 1663.66 2807.011 3613.197 5166.186 7396.243 13235.07 29705.14 121094.3

Figure 4.11: Reserves for each occurrence year from the example of the pseudo-data

We store this information and go back and re-sample with replacement our residuals in order to keep creating sets of pseudo-data and do this a total of 10,000 runs.

Year Of Occurrence	Reserve's Mean	Minimum Observed	Maximum Observed	Reserve's Standard Deviation
2008	794	-473	2,578	448
2009	1,682	-172	4,182	597
2010	2,622	687	5,105	714
2011	3,488	1,274	6,192	813
2012	4,629	2,460	7,934	856
2013	7,457	4,497	11,886	1,137
2014	13,358	9,169	18,533	1,467
2015	29,384	23,273	36,029	2,189
2016	128,198	110,119	146,458	5,465
Total	191,610	163,288	215,317	7,305

Table 4.1: Bootstrap Results

In the following plot we can find the comparison between the Normal Distribution, considering a mean of 191,610 and standard deviation of 7,305, and the empirical distribution regarding the estimator of the total reserve.

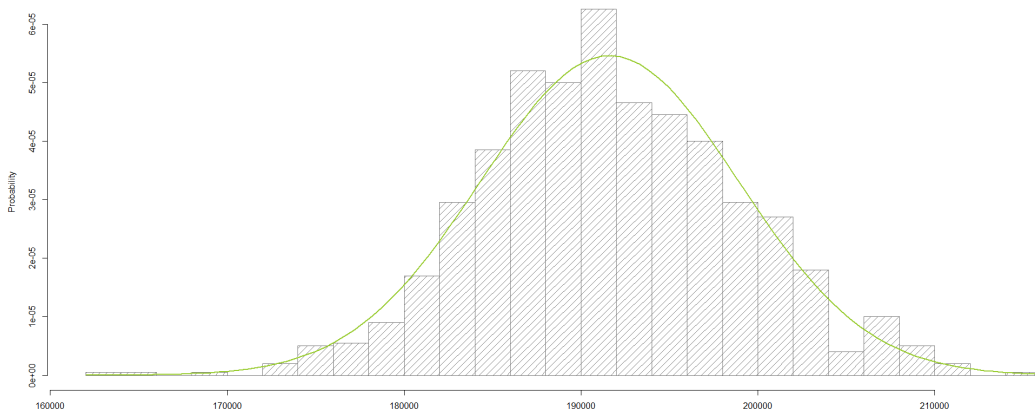


Figure 4.12: Approximate probability distribution of the Total Reserve

From Figure 4.12 we can see that the Normal distribution estimates the total reserve well enough, even with the little shift to the right. In order to prove this statement we did a Kolmogorov-Smirnov Goodness-of-Fit Test, where,

H_0 : The data follow a Normal Distribution.

vs

H_1 : The data does not follow a Normal Distribution.

Running the *ks.test* function in *R* we obtained a p-value of 0.75, meaning that we do not reject the null hypothesis.

This allows us to compute confidence intervals considering the Normal distribution.

From Section 2.4 and from the estimations above, we get the estimation of the confidence intervals given by (2.2) for each occurrence year.

	2007	2008	2009	2010	2011	2012	2013	2014	2015	2016
Lower Limit	0	-84.92876	513.024	1221.845	1894.602	2950.701	5228.828	10482.84	25093.24	117487.3
Upper Limit	0	1673.01876	2851.848	4021.485	5080.702	6306.927	9685.834	16233.88	33674.72	138908.4

Figure 4.13: Estimation of confidence interval for R_i

For the total reserve we have that for the confidence level of 95%, our confidence interval is [177293;205928]

CONCLUSION

In the last two chapters, we presented two stochastic models in order to estimate reserves. These models were applied to the same set of data, so now we are going to compare the results.

In the Table 5.1 are the values for the estimated reserves for each occurrence year and their *MSE* for both presented methods, Thomas Mack and Bootstrap. There will also be included the 95% confidence level intervals for the total reserves.

Year Of Occurrence	Thomas Mack Reserve	Bootstrap Mean Reserve
2008	800	794
	93	448
2009	1,699	1,682
	234	597
2010	2,629	2,622
	617	714
2011	3,522	3,488
	727	813
2012	4,630	4,629
	1,031	856
2013	7,467	7,457
	1,088	1,137
2014	13,313	13,358
	1,422	1,467
2015	29,371	29,384
	1,712	2,189
2016	127,789	128,198
	5,300	5,465
Total	195,286	191,610
	6,665	7,305
Lower Limit Confidence Level 95%	178,210	177,293
Upper Limit Confidence Level 95%	204,336	205,928

Table 5.1: Thomas Mack vs. Bootstrap

We can observe that both methods offer a similar estimation of reserves, which is to be expected since they both are a stochastic approach to the Chain Ladder method.

The *MSE* for the Thomas Mack is 6,665 hundreds of euros whilst for the Bootstrap

it is 7,305 hundreds of euros which represents 3% and 4%, respectively. The higher value for the Bootstrap might be due to the randomness of the residuals and the various sets pseudo-data.

Considering the 95% of confidence level in Table 5.1, we can assume from the confidence interval that a cautious value for the reserve would be 204,336 hundreds of euros. However, this value can not be taken lightly because there are other factors that might influence these estimations, minding that the actuary has to address the reality of the company as well.

The estimation of reserves is very important for the actuarial equilibrium of the non-life section for insurance companies. In order to be prudent, it is useful to apply stochastic methods, instead of deterministic methods, due to their advantages, particularly, being able to define confidence intervals so we can choose the best reserve value with high confidence.

From the Thomas Mack method we can see that these type of methods are also build with rigorous statistical foundations, which the data has to overcome in order to use tests to evaluate the adjustments made to it.

Neglecting the use of the benefits of stochastic methods might lead to estimate a reserve that could be good enough but with low certainty, putting the insurance company in risk of solvency.

Sometimes to better estimate our reserves the estimation methods must reflect the aspects of different insurance types. For this we can use a tail factor, that should be selected through the curves of adjustment of the development coefficients. In our case, this adjustment was not made due to the insignificance of the estimated tail factor.

Besides the advantages mentioned, these methods also have a low need of computational power which makes them more practical.

Finally, it is important to point out that external factors, such as legal, economical and social have to be taken into account since these are not covered by the estimations method and might alter the future reserves.

BIBLIOGRAPHY

- Autoridade de Supervisão de Seguros e Fundos de Pensão. (n.d.). <https://www.asf.com.pt>. Accessed: 2019-06-25. (Cit. on p. 4).
- England, P. D., & Verrall, R. J. (2002). Stochastic claims reserving in general insurance. *British Actuarial Journal*, 8(3), 443–518. (Cit. on pp. 23, 24).
- Gesmann, M., Murphy, D., Zhang, Y. (, Carrato, A., Wuthrich, M., Concina, F., & Dal Moro, E. (2021). *Chainladder: Statistical methods and models for claims reserving in general insurance*. R package version 0.2.14. Retrieved from <https://CRAN.R-project.org/package=ChainLadder>. (Cit. on pp. 17, 26)
- Lowe, J. (1994). A Practical Guide To Measuring Reserve Variability Using : Bootstrapping , Operational Time And A Distribution-Free Approach. *General Insurance Convention*. (Cit. on p. 24).
- Mack, T. (1993). Distribution-free Calculation of the Standard Error of Chain Ladder Reserve Estimates. *ASTIN Bulletin*, 23(2), 213–225. doi:10.2143/ast.23.2.2005092. (Cit. on pp. 9, 11, 12)
- Straub, E., & Grubbs, D. (1998). The faculty and institute of actuaries claims reserving manual. volume 1 and 2. *ASTIN Bulletin*, 28(2). doi:10.1017/S0515036100012472. (Cit. on p. 1)
- Taylor, G. (2012). *Loss reserving: An actuarial perspective*. Springer Science & Business Media. (Cit. on pp. 3, 4, 23).



