

ON THE PREDICTION ERROR IN SEVERAL
CLAIMS RESERVES ESTIMATION METHODS

IEVA GEDIMINAITĖ

MASTER THESIS
STOCKHOLM, 2009

Royal Institute of Technology
School of Engineering Sciences
Dept. of Mathematics
Master Thesis in Insurance Mathematics

ON THE PREDICTION ERROR IN SEVERAL CLAIMS RESERVES ESTIMATION METHODS

Ieva Gediminaitė

Supervisors: Professor Boualem Djehiche

Chief Actuary Lars Klingberg

2009

ABSTRACT

The Chain-Ladder (CL) and the Bornhuetter-Ferguson (BF) reserve estimation methods are the most common in the general insurance reserving process. It is very important to know how accurate the resulting estimates are. Mack derived theoretical prediction error formulae for the CL method in 1993 and for the BF method in 2008. Also, bootstrap technique for the CL method has been introduced and developed since England & Verrall published their work in 1999.

In this thesis the theory behind all the calculations is first explained. Then, the personal accident (PA) insurance data from Trygg-Hansa Försäkrings AB is analyzed. First, the theoretical prediction error for both methods is calculated, according to the articles by Mack. Second, the most recently developed bootstrap procedure is applied for the CL method. Finally, a bootstrap procedure to calculate the estimation error of the BF method is constructed and applied on the PA data. A comparison study of all the performed computations is then given.

Key words and phrases: Prediction error, Chain-Ladder, Bornhuetter-Ferguson, Stochastic claims reserving, Bootstrap in claims reserving.

ACKNOWLEDGEMENTS

I thank all that have been helpful during the time of writing this thesis. In particular, thanks to my supervisor professor Boualem Djehiche at KTH for guidance and help. Also, I thank my advisor chief actuary Lars Klingberg at Trygg-Hansa AB for opening the door to the uncertainty world and sharing his valuable experience. I am very grateful to my colleagues at Trygg-Hansa AB who helped dealing with the data and understanding tricky bootstrap techniques. A special thanks to Susanna Björkwall for her valuable comments. The last but not the least, thanks to my husband and family for all the energy and happiness they bring me.

CONTENTS

1. INTRODUCTION	7
2. PREPARATORY CONCEPTS	8
2.1 RUN-OFF TRIANGLE	8
2.2 CLAIMS RESERVING	9
2.3 BOOTSTRAP TECHNIQUE	11
3. RESERVING METHODS	13
3.1 THE CHAIN-LADDER METHOD	13
3.1.1 STOCHASTIC MODEL	13
3.1.2 CHECK OF ASSUMPTIONS	15
3.1.3 CALCULATING THE PREDICTION ERROR	18
3.1.4 BOOTSTRAPPING THE PREDICTION ERROR	19
3.2 THE BORNHUETTER-FERGUSON METHOD	26
3.2.1 STOCHASTIC MODEL	27
3.2.2 PARAMETER ESTIMATION	28
3.2.3 CALCULATING THE PREDICTION ERROR	30
3.2.4 BOOTSTRAPPING THE ESTIMATION ERROR	31
4. DATA ANALYSIS	34
4.1 PRODUCT PRESENTATION	34
4.2 RESULTS	35
4.2.1 THE CHAIN-LADDER RESULTS	36
4.2.1.1 Stochastic Assumptions	36
4.2.1.2 Theoretical Prediction Error	36
4.2.1.3 Bootstrapped Prediction Error	39
4.2.2 THE BORNHUETTER-FERGUSON RESULT	44
4.2.2.1 Estimation of Prior Ultimate Claims	44
4.2.2.2 Estimation of other BF parameters	49
4.2.2.3 Theoretical Prediction Error	52

4.2.2.4	Bootstrapped Estimation Error	58
4.2.3	CHAIN-LADDER VS BORNHUETTER-FERGUSON	61
4.2.3.1	Theoretical Prediction Error	61
4.2.3.2	Bootstrapped Estimation Error	65
5.	CONCLUSIONS	68
6.	REFERENCES	69
7.	APPENDIX	72
7.1	APPENDIX I. THE CL PREDICTION ERROR	72
7.2	APPENDIX II. THE BF PREDICTION ERROR	78
7.3	APPENDIX III. (CL1) AND (CL3) CHECK	82
7.3.1	APPENDIX III.A – INCURRED CUMULATIVE CLAIMS LINEARITY (CL1) CHECK	82
7.3.2	APPENDIX III.B – PAID CUMULATIVE CLAIMS LINEARITY (CL1) CHECK	85
7.3.3	APPENDIX III.C – INCURRED CLAIMS DATA VARIANCE ASSUMPTION (CL3)	88
7.3.4	APPENDIX III.D – PAID CLAIMS DATA VARIANCE ASSUMPTION (CL3)	92
7.4	APPENDIX IV. THE ESTIMATION ERROR OF THE BF METHOD	96

1. Introduction

Reserving actuaries in their daily work are trying to predict the future – what reserve amount is sufficient to pay the clients given that the future payments can vary. Trygg-Hansa is one of the largest general insurance companies in Sweden. Its whole variety of products is divided into several insurance classes with varying payment patterns. In particular, Personal Accident (PA) insurance classes often have late claim development which encourages actuaries to try a variety of prediction models. For example, in addition to the widely used Chain-Ladder (CL) method, which is build on historical claim development, it is recommended to also work with the Bornhuetter-Ferguson (BF) method which is built on an exposure measure. The combination of the two methods provides more reliable reserve estimates for this class.

In the process of estimation it is very important to realize how accurate the estimated results are. The prediction error of the reserve estimate for each accident period as well as for the overall result shows how widely spread the future values might be. Trygg-Hansa actuaries are using the EMB ResQ Professional reserving tool for their analysis. In the software one can find standard error computed for the CL method. The BF method in ResQ provides actuaries with the estimated reserve but not with the standard error of the estimates. Analyzing PA insurance data actuaries estimate the reserve by combining the two methods. This is why it is important to implement the calculation of the prediction error for the BF method.

The goal of this thesis is to compute the prediction error for the CL and the BF reserve estimates from already derived formulae; also, to apply the existing bootstrap technique for the CL method and to construct a similar bootstrap technique for the BF method. In the section 2 basic concepts used throughout the work will be introduced. In the section 3 the two reserving methods will be presented and the theoretical calculation of the prediction error explained. The bootstrap procedures for both methods are outlined in the corresponding subsections of the section 3, data analysis with results and comparison study follows in the section 4. The main concluding statements and further possible investigation is given in the section 5. In the appendix the most important proofs as well as some extra graphs for deeper understanding are attached.

2. Preparatory Concepts

2.1 Run-off Triangle

Before the reserving methods are discussed, the term *run-off triangle* must be presented. The run-off triangle is a convenient way to present observed data. It is an upper-triangle with accident period values in its rows and development period in its columns. In the following text the data will be analyzed quarterly, therefore both accident and development periods are quarters.

An insurance company sells different kinds of insurances. Each contract has its own number with complementary variables showing the kind of insurance and other conditions of the agreement. Let us analyze one particular agreement A. Say, it was signed in 2008-01-14, and it is a personal accident insurance. In June, 2008, the insured person was camping and accidentally broke his arm when climbing a tree on 2008-06-20. After he got home, in 2008-06-27, he called the insurance company and claimed for a proper sum of money, say M, that would cover hospital and other costs. After the paper work was finished, the insured person received the money M from the insurance company on 2008-10-03. In the insurance company's database this would be depicted as follows:

Accident date: 2008-06-20

Registration date: 2008-06-27

Payment date: 2008-10-03

Now, say, we are at January 2009, and we want to place this claim into a proper place of run-off triangle for data analysis. The claim happened on 2008-06-20, therefore accident quarter is 2008Q2. The payment was done, i.e. the claim developed, on 2008-10-03, therefore development quarter is 2008Q4, or 3rd development quarter for accident period 2008Q2. In the Figure 1 you can see where the amount of money paid, M, should be added.

Dev qtr Acc qtr	1	2	3	4	5	6
...	...	...	...	...	...	...
...	...	...	...	...	...	
...	...	...	...	...		
2008Q2	...	...	... + M			
2008Q3	...	...				
2008Q4	...					

Figure 1. The Run-off triangle.

It is intuitively clear that an insurance company will get thousands of claims every quarter, and the run-off triangle is just the way how all these claims are summarized for the further analysis.

2.2 Claims Reserving

Considering the previous example, where the claim was closed so quickly (after 3,5 months), and the amount paid (M) was equal to the amount which the client claimed for, it might seem that employing actuaries to compute reserves is almost unnecessary. Here one should beware that not all the claims, or, very little part of the claims history end up so clearly and quickly as in the previous example. Depending on the insurance policy, claims to the insurance company might come after several years an accident has occurred, or, the client might claim for one sum of money, but in the end it appears that the company has to pay much more, or much less. All these scenarios must be predicted as good as possible, and this is where claims reserving methods come into hand.

Recall the run-off triangle with the data observed (Figure 1). We do not know yet how many and how big claims will be in the cell for accident period 2008Q4 in the 2nd development period, or for accident period 2008Q2 in the 5th development period. Moreover, we do not know any of the figures in the lower triangle (red/colored cells in Figure 2). These numbers are random, and a good stochastic model is needed to make precise expectations about the future. The known claims development, the run-off triangle, is usually used for drawing conclusions about the future.

Dev qtr \ Acc qtr	1	2	3	4	5	6
...	...	...	...	...	...	...
...	...	...	...	...	...	...
...	...	...	...	...	...	...
2008Q2	...	...	... + M	...	...	...
2008Q3	...	...	...	...	...	...
2008Q4	...	...	...	...	...	...

Figure 2. The lower triangle of the expected claims development in the future.

Assume that the run-off triangle is filled with the incremental claim amounts, then the future incremental claims (to be filled in the red/colored cells) is what a company did not yet register, but expects to register. It is the so called *claims reserve*. More common is to have a cumulative claims run-off triangle. In this case, the reserve is the difference between the figures in the last development period (in Figure 2 it is column 6) and the latest known cumulative claims (the last white color diagonal). Notice also that we are trying to fill in future claims only for those accident periods that already have passed, i.e. the claims we are predicting have already occurred (incurred) but were not reported. Due to this definition the term IBNR (**I**ncurred **B**ut **N**ot **R**eported) was

created. Various companies interpret the definition differently. Sometimes the IBNR might include not only reserve for the claims that have not yet been reported, but also those claims that have been not enough reported, i.e. it is believed that something more will happen with those claims. To conclude, the prediction of ultimate claims, the expectation of the level when all claims for a particular accident period are settled, is one of the most important tasks for a reserving actuary to make. As R. L. Bornhuetter and R. E. Ferguson indicate in their work *The Actuary and IBNR* (Bornhuetter&Ferguson, 1972), the IBNR reserve calculation is not only important for the company but is required by law. It is important to understand that inaccurate IBNR reserves will lead to non-optimal management decisions.

Quite a few methods were created so far for predicting the future claims. To mention some of them: the Chain-Ladder method, the Naïve loss ratio method, the Bornhuetter-Ferguson method, the Cape-Cod method, the Benktander method. All these methods use historical claims development in one way or another, and provide actuaries only with the point estimate of the reserve. It is obvious that in order to choose good estimates actuaries might reach for knowledge about any statistical inference of the estimates, for example variability of them. To find out the variability of an estimate first, one has to have a stochastic model for the data and second, to understand the distribution of the estimator. Most investigation has been done on the two commonly used reserving methods, the Chain-Ladder (CL) method and the Bornhuetter-Ferguson (BF) method. In 1993 T. Mack has derived the formula for computing prediction error (variability) for the CL reserve estimate (Mack, 1993). The formula is very widely used and helps actuaries understanding the risk and responsibility of predicting the future. The investigation of the BF method estimates and derivation of the prediction error formula was published only a year ago in Mack (2008). This formula has not yet spread out as widely as the one for the CL method, but knowing the importance of the problem, the use of it should be implemented as soon as possible.

In the next section the two reserving methods, the CL and the BF, will be introduced, meanwhile some commonly used notations (Mack, 2008) will be presented:

$C_{i,k}$ – cumulative claims amount of accident period i after k development periods

v_i - premium volume of accident period i

n - most recent accident period

$S_{i,k} = C_{i,k} - C_{i,k-1}$ - incremental claims amount of accident period i

U_i - the ultimate claims amount

$R_i = U_i - C_{i,n+1-i}$ - claims reserve for accident period i

$S_{i,n+1} = U_i - C_{i,n}$ - incremental claims amount after development period n (tail development)

F_{ij} - the factor that would develop losses from development period j to the end for accident period i

L_i - claims relative to an exposure (ultimate loss ratio for accident period i)

f_j – development factor for development period j

z_j – the estimated percentage of the ultimate claims amount that is expected to be known after development period k

σ_j^2, s_j^2 - proportionality constants for development period j

x_i, y_i - unknown parameters

In the above, cumulative corresponds to claim amounts (paid or incurred) that had been registered *during and up to* the period, while incremental amount is the registered amount *during* the period. The ultimate claims amount is the total amount of claims for accident period i . It is expected that all claims' development is included in the ultimate amount. The latest run-off triangle diagonal is the current claims amount, and subtracting it from the ultimate claims gives the reserve. In other words, the reserve is the expectation of further development of claims. Finally, the ultimate loss ratio is one of the key ratios in actuarial mathematics. It is calculated as ultimate claims divided by premiums, and it shows how good the business is. The lower the ultimate loss ratio, the better the business is.

Having presented the notations and concepts from the actuarial field, let us proceed with introducing the bootstrap technique.

2.3 Bootstrap Technique

Statisticians often find themselves in situations when they have gathered data and they need to estimate some parameters from the data, preferably not only the point estimates, but also confidence intervals and as many as possible other statistics. The task is rather difficult when one does not know the distribution of the parameter. Moreover, it might seem almost impossible, if one does not even know how the parameter is calculated, he/she is just having a hard-coded function calculating the parameter from the data. In the end of the 20th century a technique making all this

possible was developed. It is called *bootstrapping*, and, as one could expect, it requires powerful computers to implement statistical inferences about the parameters.

A very good theoretical background as well as plenty examples of bootstrap technique can be found in the book *An Introduction to the Bootstrap* by Efron & Tibshirani (1994). Bootstrap technique is based on making a copy of the original experiment. One can choose whether to use parametric or non-parametric bootstrap. In the parametric bootstrap a parametric model for the data should be chosen, parameters should be estimated, and then new datasets should be generated. This method depends heavily on the chosen parametric model. Instead, one could choose to use non-parametric bootstrap. In that case the empirical distribution $\hat{F}_n$ is used as an estimate of F . The new datasets are generated from the empirical distribution, which is done by random sampling with replacement from the original dataset.

In this work the non-parametric bootstrap technique will be used on incremental claims data to obtain a distribution of the prediction error. Since the incremental claims data form a linear model, the bootstrap procedure has to be adjusted to that. In the theory (Efron&Tibshirani, 1994) two ways of bootstrapping in the linear model are presented:

- Paired bootstrap, or bootstrap of points – the resampling is done directly from the observations.
- Bootstrap of residuals – the resampling is applied to the residuals of the model.

Which of the two ways should be chosen depends a lot on the situation. In general, the paired bootstrap is more robust than the residuals bootstrap. Moreover, the residuals are not necessarily outcomes from independent and identically distributed random variables. Most common is that the residuals are larger in the central areas, and smaller in the tails. Despite all the disadvantages, the residuals bootstrap only can be used in claims reserving, given the dependence between some observations and the parameter estimates. Therefore, a good way of standardizing residuals will have to be found to be able to assume the independence. To conclude, assuming the computed residuals are outcomes from independent and identically distributed random variables, we can bootstrap the residuals, and generate a new sample of observations by calculating backwards in the residuals' formula. The latter approach will be used in bootstrapping the Chain-Ladder prediction error estimator, as well as the Bornhuetter-Ferguson one.

In the following section both reserving methods will be presented, and formulae as well as bootstrap procedures for the methods will be outlined.

3. Reserving Methods

3.1 The Chain-Ladder Method

The most common method for claims reserving is the Chain-Ladder (CL) method. It does not require any advanced programs, one can implement the method in any spreadsheet, therefore, it is simple to use. One more advantage of the method is that it is possible to fit a distribution-free stochastic model to the method (Mack, 1993).

Much investigation has been made on the CL method. As mentioned above, it is very important to have not only a reserve estimate but also the volatility of the estimate. Many papers have been published with the same attempt – to fit the best stochastic model for the CL method. With a proper stochastic model one can estimate the prediction error of an estimate fairly easily. Already Hachemeister & Stanard (1975) offered Poisson distributed incremental claims stochastic model which led to the estimates very close to the original CL estimates. A least squares regression approach was used in a few papers (Taylor & Ashe (1983), Zehnwirth (1989), Renshaw (1989), Christofides (1990), Verrall (1991)). Wright (1990) tried the generalized linear model and the method of scoring, and gamma distribution and maximum likelihood estimation was offered by Mack (1991). A distribution-free formula for the standard error of the CL reserve estimates was presented by Mack (1993). The author was awarded as a joint winner in the Casualty Actuarial Society (CAS) prize paper competition on variability of loss reserves. The latter, distribution-free, model is used as a base when computing variability of the estimate in the previously mentioned program EMB ResQ Professional, and was also chosen to be analyzed in this work with personal accident claims data.

One may argue if the so called “distribution-free” model is really free of distribution, since it still has quite strong assumptions on the raw data. In the following section the assumptions (Mack, 1993) will be presented.

3.1.1 Stochastic Model

The first assumption indicates that the data from one development period to the other should construct a linear regression with the factor f_j . The factors are called development factors, link ratios, or age-to-age factors. They are the core estimation for the CL method. As can be seen from the formula (1) below, they are computed in a very simple way. In the Figure 3 one can see the first, third and fifth development factors’ construction. On the other hand, one should look at it as a weighted average of development factors in different accident periods for the same development period:

CL1 $E(C_{i,j+1}|C_{i,1}, \dots, C_{i,j}) = C_{i,j} * f_j$, with

$$\hat{f}_j = \frac{\sum_{i=1}^{n-j} C_{i,j+1}}{\sum_{i=1}^{n-j} C_{i,j}} = \sum_{i=1}^{n-j} \frac{C_{i,j}}{\sum_{i=1}^{n-j} C_{i,j}} \cdot \frac{C_{i,j+1}}{C_{i,j}} \quad (1)$$

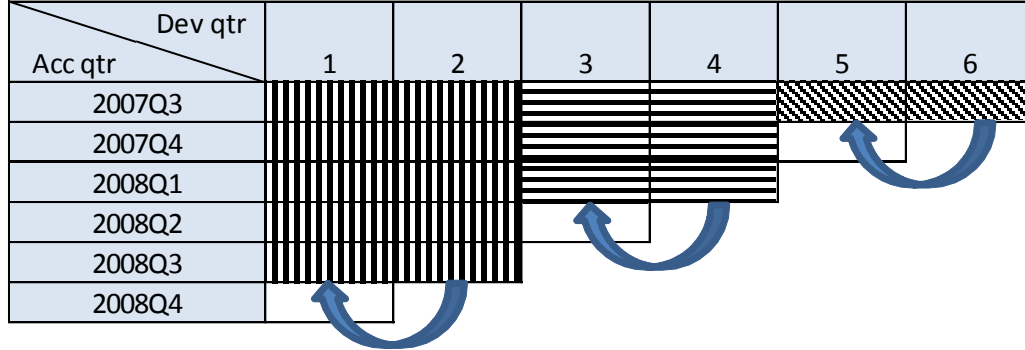


Figure 3. The construction of the 1st, 3rd, and 5th development factors.

Secondly, the distribution-free model assumes that cumulative claims in one accident period are independent from cumulative claims in another accident period.

CL2 Vectors $\{C_{i,1}, \dots, C_{i,j}\}$ and $\{C_{k,1}, \dots, C_{k,j}\}$ are independent $\forall i \neq k$.

This is a strong assumption and practically one must be very careful about it. In real life there might be changes in claims handling or case reserving internally, as well as changes in court decisions and inflation externally. This would create dependencies among different accident quarters. In this case the actuary should make himself/herself aware of that and adjust the reserving methodology to the events in the market.

The final assumption of the method is about the variability of the raw data. It is assumed that there exists a proportionality factor σ_j^2 , which relates variance of cumulative claim amount with a cumulative claim amount one development period earlier.

CL3 $Var(C_{i,j+1}|C_{i,1}, \dots, C_{i,j}) = C_{i,j} * \sigma_j^2$, where σ_j^2 is unknown parameter, estimated by

$$\hat{\sigma}_k^2 = \frac{1}{n-k-1} \sum_{i=1}^{n-k} C_{i,k} \left(\frac{C_{i,k+1}}{C_{i,k}} - \hat{f}_k \right)^2, 1 \leq k \leq n-2 \quad (2)$$

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

Formula (3) is needed because there is too little data to calculate standard deviation for the latest period in a regular way. The $\hat{\sigma}_{n-1}^2$ can be extrapolated using the calculated series, or, like Mack (1993) suggested, it can be calculated requiring that

$$\frac{\hat{\sigma}_{n-3}}{\hat{\sigma}_{n-2}} = \frac{\hat{\sigma}_{n-2}}{\hat{\sigma}_{n-1}}, \text{ if } \hat{\sigma}_{n-3} > \hat{\sigma}_{n-2}$$

which leads to the (3) formula (Mack, 1993).

Notice, under (CL1) and (CL2) the estimators $\hat{f}_j, 1 \leq k \leq n-1$, are unbiased and uncorrelated. The proof can be found in the Appendix I, or Mack (1993).

3.1.2 Check of Assumptions

There is a way to check if the data to be analyzed fulfils the three assumptions. The procedures are taken from Mack (1994) and will be described here, also implemented for the chosen data later.

First, we have to check if there is a linear relation between data in two adjacent development periods. The values of $C_{i,k}, 1 \leq i \leq n, k \text{ fixed}$, are to be considered as non-random values and equations in (CL1) can be interpreted as an ordinary regression model of the type

$$Y_i = c + x_i b + \epsilon_i, 1 \leq i \leq n$$

where c and b are regression coefficients and ϵ_i is the error term with $E(\epsilon_i) = 0$. In the CL assumption case $c = 0, b = f_k$, and $Y_i = C_{i,k+1}$ at the points $x_i = C_{i,k}$ for $i = 1, \dots, n-k$. The regression coefficient $b = f_k$ could be estimated by the usual least squares method, but in that case we would not get an estimator for f_k which would be in line with later variance assumption. Mack (1994) indicates that one should bear in mind that the least squares method implicitly assumes equal variances $Var(Y_i) = Var(\epsilon_i) = \sigma^2, \forall i$. If variances do depend on accident period, one should use a weighted least squares approach which consists of minimizing the weighted sum of squares

$$\sum_{i=1}^n w_i (Y_i - c - x_i b)^2$$

where the weights w_i are in inverse proportion to $Var(Y_i)$. In order to agree with (CL3) regression weights proportional to $1/(C_{i,k} \alpha_k^2)$ should be used. Then the following should be minimized:

$$\sum_{i=1}^{n-k} (C_{i,k+1} - C_{i,k} f_k)^2 / (C_{i,k} \alpha_k^2)$$

Since k is fixed, α_k^2 does not add to the minimization problem, therefore, minimizing

$$\sum_{i=1}^{n-k} (C_{i,k+1} - C_{i,k} f_k)^2 / C_{i,k}$$

with respect to f_k yields to the CL estimator $\hat{f}_k$.

After the estimation of the development factors into the regression framework, the usual regression analysis instruments are available to check the underlying assumptions of linearity and the variance. Plotting $C_{i,k+1}$ against $C_{i,k}$, $i = 1, \dots, n - k$, linear relationship between data in two development periods can be checked. If the points are located around a straight line through the origin with slope $\hat{f}_k$, the (CL1) assumption is accepted. Secondly, if linearity was accepted, the weighted residuals

$$\frac{C_{i,k+1} - C_{i,k} \hat{f}_k}{\sqrt{C_{i,k}}}, 1 \leq i \leq n - k$$

should be plotted against $C_{i,k}$. To accept the variance assumption the plot of residuals should not show any particular trends but appear purely random.

We can now turn to investigating the independence of accident quarters. The reasons which could disturb the independence were mentioned above. It is important to realize, that the effect in the data would be seen on calendar quarters (diagonally). Mack (1994) offered the following testing procedure.

A calendar period influence affects one of the diagonals

$$D_j = \{C_{j,1}, C_{j-1,2}, \dots, C_{2,j-1}, C_{1,j}\}, 1 \leq j \leq n$$

and therefore also influences the adjacent development factors

$$A_j = \left\{ \frac{C_{j,2}}{C_{j,1}}, \frac{C_{j-1,3}}{C_{j-1,2}}, \dots, \frac{C_{1,j+1}}{C_{1,j}} \right\}$$

and

$$A_{j-1} = \left\{ \frac{C_{j-1,2}}{C_{j-1,1}}, \frac{C_{j-2,3}}{C_{j-2,2}}, \dots, \frac{C_{1,j}}{C_{1,j-1}} \right\}$$

where the elements of D_j form either the denominator or the numerator. Thus, if due to a calendar period influence the elements of D_j are larger (smaller) than usual, then the elements of A_{j-1} are also larger (smaller) than usual and the elements of A_j are smaller (larger) than usual.

Therefore, in order to check for such calendar quarter influences we only have to subdivide all development factors into ‘smaller’ and ‘larger’ ones and then to examine whether there are diagonals where the small development factors or the large ones clearly prevail. For this purpose we divide the column of all development factors observed between development quarters k and $k + 1$ into two groups LF_k of larger factors being greater than the median of the column, and SF_k of smaller factors below the median of the column. If the number of elements $n - k$ in a column is odd there is one element which is equal to the median of the column and is not assigned to any of the two groups.

Having done this procedure for each column in the run-off triangle every development factor observed is

- either eliminated (as equal to the median)
- or assigned to the set of larger factors L
- or assigned to the set of smaller factors S

In this way every development factor which is not eliminated has a 50% chance of belonging to either L or S .

Now we count for every diagonal $A_j, 1 \leq j \leq n - 1$ of development factors the number L_j of large factors and the number S_j of small factors. Intuitively, if there is no specific change from calendar period j to calendar period $j + 1$, A_j should have about the same number of small factors as of large factors, i.e. L_j and S_j should be of approximately same size apart from pure random fluctuations. But if L_j is significantly larger or smaller than S_j or, equivalently, if $Z_j = \min(L_j, S_j)$ is significantly smaller than $\frac{L_j + S_j}{2}$, then there is some reason for a specific calendar period influence.

A formal test was designed in Mack (1994) and here we give only the resulting formulae. Consider

$$Z = Z_2 + \dots + Z_{n-1}$$

Since under null-hypothesis different Z_j 's are (almost) uncorrelated we have

$$E(Z) = E(Z_2) + \dots + E(Z_{n-1})$$

$$Var(Z) = Var(Z_2) + \dots + Var(Z_{n-1})$$

It can be shown (Mack, 1994) that

$$E(Z_j) = \frac{n}{2} - \binom{n-1}{m} \frac{n}{2^n}$$

$$\text{Var}(Z_j) = \frac{n(n-1)}{4} - \binom{n-1}{m} \frac{n(n-1)}{2^n} + E(Z_j) - (E(Z_j))^2$$

with $n = L_j + S_j$ and $m = \lfloor (n-1)/2 \rfloor$.

We can assume that Z has approximately (check with the article) a normal distribution. This means that we reject (with an error probability of 5%) the hypothesis of having no significant calendar period effects only if not

$$E(Z) - 2\sqrt{\text{Var}(Z)} \leq Z \leq E(Z) + 2\sqrt{\text{Var}(Z)}$$

3.1.3 Calculating the Prediction Error

Having stochastic model assumptions for the CL method stated and checked, the formula for computing the prediction error of the reserve estimate for each accident quarter as well as total will be given. All formulae presented in this section were derived by Mack (1993). In the Appendix I the proofs (Mack, 1993) with minor clarifications added can be found.

Before stating the main result some important things have to be observed. When evaluating the mean square error we will not be using the unconditional mean squared error. Instead, we are more interested in the conditional mean squared error of the particular estimated amount $\widehat{C}_{i,n}$ based on the specific data set D observed. This will just give us the average deviation between $\widehat{C}_{i,n}$ and $C_{i,n}$ due to future randomness only. Therefore, the mean squared error $mse(\widehat{C}_{i,n})$ we want to estimate is defined to be $mse(\widehat{C}_{i,n}) = E((\widehat{C}_{i,n} - C_{i,n})^2 | D)$, where $D = \{C_{i,k} | i+k \leq n+1\}$ is the set of all data observed so far.

First, notice

$$mse(\widehat{R}_i) = E((\widehat{R}_i - R_i)^2 | D) = E((\widehat{C}_{i,n} - C_{i,n})^2 | D) = mse(\widehat{C}_{i,n})$$

Next, because of the general rule $E(X - a)^2 = \text{Var}(X) + (E(X) - a)^2$ we have

$$mse(\widehat{C}_{i,n}) = \text{Var}(C_{i,n} | D) + (E(C_{i,n} | D) - \widehat{C}_{i,n})^2 \quad (4)$$

which shows that the mean square error is the sum of two terms – the stochastic error (process variance) $\text{Var}(C_{i,n} | D)$ and the estimation error $(E(C_{i,n} | D) - \widehat{C}_{i,n})^2$.

Theorem 1. Under the assumptions (CL1), (CL2), and (CL3) the mean square error $mse(\widehat{R}_i)$ can be estimated by

$$\widehat{mse}(\hat{R}_l) = \hat{C}_{l,n}^2 \sum_{k=n+1-i}^{n-1} \frac{\hat{\sigma}_k^2}{\hat{f}_k^2} \left(\frac{1}{\hat{C}_{l,k}} + \frac{1}{\sum_{j=1}^{n-k} C_{j,k}} \right) \quad (5)$$

where $\hat{C}_{l,k} = C_{i,n+1-i} \hat{f}_{n+1-i} \dots \hat{f}_{k-1}$, $k > n+1-i$ are the estimated values of the future $C_{l,k}$ and $\widehat{C}_{l,n+1-i} = C_{l,n+1-i}$.

The standard error, or as called earlier, the prediction error of the reserve estimator is the square root of the mean squared error:

$$s.e.(\hat{R}_l) = \sqrt{\widehat{mse}(\hat{R}_l)}$$

So far a formula for calculating the prediction error of the reserve estimate for each accident quarter was presented. It is often the case, that one is interested in the overall reserve estimate and its variability. The overall reserve estimate by itself is very simple:

$$\hat{R} = \hat{R}_2 + \dots + \hat{R}_n$$

To compute the standard error for the overall reserve estimate we have to take into account that $\hat{R}_i$'s are correlated via the common estimators $\hat{f}_k$ and $\hat{\sigma}_k$.

Corollary 1. With the assumptions and notations of Theorem 1 the mean squared error of the overall reserve estimate $\hat{R} = \hat{R}_2 + \dots + \hat{R}_n$ can be estimated by

$$\widehat{mse}(\hat{R}) = \sum_{i=2}^n \left\{ \left(s.e.(\hat{R}_i) \right)^2 + \hat{C}_{i,n} \left(\sum_{j=i+1}^n \hat{C}_{j,n} \right) \sum_{k=n+1-i}^{n-1} \frac{\frac{2\hat{\sigma}_k^2}{\hat{f}_k^2}}{\sum_{l=1}^{n-k} C_{l,k}} \right\} \quad (6)$$

Theoretical estimates are good to have, but so far we cannot conclude anything about the distribution of the prediction error. The bootstrap technique will help us to answer some more questions about the CL prediction error.

3.1.4 Bootstrapping the Prediction Error

There have been several investigations made concerning the estimation of the prediction error for the CL method. In 1999 the article *Analytic and bootstrap estimates of prediction errors in claims reserving* by P. England and R. Verrall was published (England & Verrall, 1999). The authors discuss one possible way to implement bootstrap technique and then compare it with the parametric models' estimates of the prediction error. Another approach to use the bootstrap methodology in estimating the prediction error of the CL method was presented in the article *Bootstrap Methodology*

in *Claims Reserving* written by Paulo J.R. Pinheiro, João M. Andrade e Silva, and Maria de Lourdes Centeno in 2001 (Pinheiro et al., 2001). The main difference between the two bootstrapping techniques is estimation of the process error. If the one presented by England & Verrall (1999) is applicable only for plain chain-ladder technique, the Pinheiro et al (2001) technique allows for adjustments to any development factors' method. The difference is explained more thoroughly by Susanna Björkwall in her licenciante thesis (Björkwall, 2009). The generalized version of bootstrap technique from the thesis (Björkwall, 2009) was borrowed to construct the following procedure:

- I. The preliminaries
 - Estimate the development factors $f_1, f_2, \dots, f_{n-1}$ for the triangle of the data as in (1)
 - Calculate the fitted values $\hat{m}_{ij}, i = 1, 2, \dots, n, j = 1, 2, \dots, n + 1 - i$ and the future expected values $\hat{m}_{ij}, i = 1, 2, \dots, n, j = n + 2 - i, \dots, n$
 - Calculate the residuals
 - Calculate the outstanding claims $\hat{R}_i = \sum_{j=1}^{n+1-i} \hat{m}_{ij}$ and $\hat{R} = \sum_{i=1}^n \sum_{j=1}^{n+1-i} \hat{m}_{ij}$
- II. Bootstrap world (to be repeated B times)
 - i. The estimated outstanding claims
 - Resample the residuals obtained in the first stage using replacement
 - Create pseudo-data by solving the residuals' formula backwards with $\hat{m}_{ij}, i = 1, 2, \dots, n, j = 1, 2, \dots, n + 1 - i$
 - Estimate the development factors with the pseudo-data and obtain the bootstrap forecast $\hat{m}_{ij}^*, i = 1, 2, \dots, n, j = n + 2 - i, \dots, n$
 - Calculate the estimated outstanding claims $\hat{R}_i^* = \sum_{j=1}^{n+1-i} \hat{m}_{ij}^*$ and $\hat{R}^* = \sum_{i=1}^n \sum_{j=1}^{n+1-i} \hat{m}_{ij}^*$
 - ii. The true outstanding claims
 - Resample again the residuals obtained in the first stage and select with replacement
 - Create pseudo-reality C_{ij}^{**} by solving the residuals' formula backwards with $\hat{m}_{ij}, i = 1, 2, \dots, n, j = n + 2 - i, \dots, n$
 - Calculate the true outstanding claims $R_i^{**} = \sum_{j=1}^{n+1-i} C_{ij}^{**}$ and $R^{**} = \sum_{i=1}^n \sum_{j=1}^{n+1-i} C_{ij}^{**}$
 - Store the prediction errors $pe_i^{**} = R_i^{**} - \hat{R}_i^*$ and $pe^{**} = R^{**} - \hat{R}^*$
- III. Bootstrap data analysis
 - Obtain the predictive distribution of R_i and R , the true outstanding claims in the real world, by plotting the B values of $\tilde{R}_i^{**} = \hat{R}_i^* + pe_i^{**}$ and $\tilde{R}^{**} = \hat{R}^* + pe^{**}$

The procedure is shown visually in the Figures 4, 5, and 6 below. Some steps in the procedure need an extra discussion. First, it is important to notice, that the residuals are calculated for incremental claims data. Having the cumulative claims, the development factors are calculated. The fitting step is

made by taking the latest diagonal of cumulative data triangle and computing backwards using the estimated development factors. Now, having the fitted cumulative claims we have to calculate the incremental claims, and only then compute residuals between the fitted incremental claims and the original incremental claims. This manipulating of data is clearly depicted in the Figure 4 below.

The second subject to discuss is how the residuals are computed. Mostly Pearson residuals are used, but various authors suggest various ways of standardizing or scaling the residuals. Pinheiro et al. (2001) investigation offers using Pearson residuals:

$$r_{ij}^{(P*)} = \frac{y_{ij} - \hat{m}_{ij}}{\sqrt{\widehat{Var}(m_{ij})}}$$

standardizing it with the factor $\frac{1}{\sqrt{1-h_{ij}}}$, where h_{ij} is the corresponding element of the diagonal of the “hat” matrix (Clarke, 2008). On the other hand, England & Verrall (1999) suggested using unscaled Pearson residuals:

$$r_p = \frac{y_{ij} - \hat{m}_{ij}}{\sqrt{\hat{m}_{ij}^p}}$$

with a global adjusting factor $\sqrt{\frac{n}{n-q}}$, where n stands for the number of observations and q – the number of parameters to estimate.

The Capital Modelling team in Trygg-Hansa are using EMB Igloo Professional software for bootstrap simulations, where residuals are calculated as unscaled Pearson with a chosen parameter p and a certain adjustment for negative incremental claims. Further the calculated residuals are standardized with chosen coefficients. The exact formulas for these calculations cannot be given due to confidentiality reasons.

To choose the best residuals’ formula various ways were tried and the one giving the most independent and identically distributed residuals for the data was chosen.

When using the England & Verrall (1999) suggested residuals calculation with $p = 2$ rather dependent residuals were obtained. In the Figures 7 and 8 below the plotted residuals are presented. If paid claims residuals could be thought of as independent, the incurred claims residuals are clearly dependent, therefore a better residuals formula should be used. The best result was received using the formula from EMB Igloo Professional. The plotted residuals are presented in the Figure 9 for paid claims data and Figure 10 for incurred claims data. The graphs allow us to assume independent and

identically distributed residuals. Therefore, using the Trygg-Hansa Capital team's formula we can implement the bootstrap procedure for the CL method.

Even though the CL method is widely used in practice, the method has some disadvantages also. As Mack (2008) indicates, the CL reserve is directly proportional to the claims amount known so far, and it only considers the development until a given last development period (no tail development). As for prediction error, the results might be very volatile especially for the latest accident periods due to too little data observed. To deal with the disadvantages of the CL method Bornhuetter & Ferguson (1972) introduced a different method for claims reserving. The method is called by the two scientists' names: the Bornhuetter-Ferguson (BF) method.

Bootstrap technique for the CL method

I. The Preliminaries

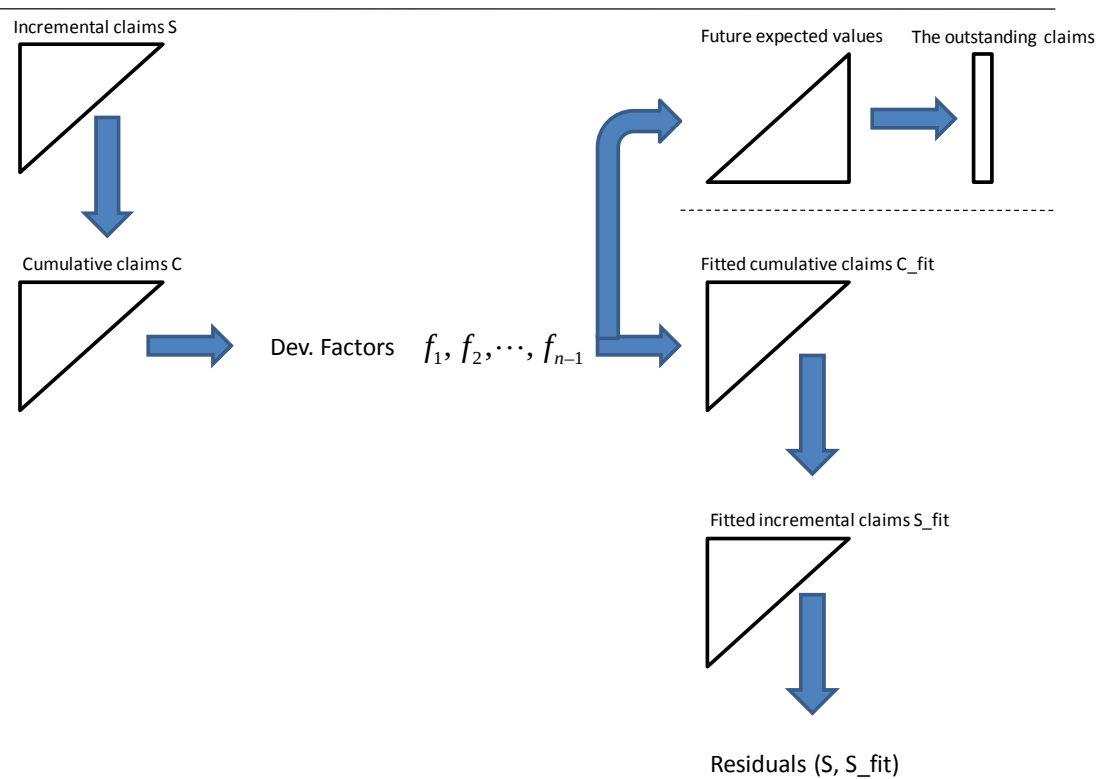


Figure 4

Bootstrap technique for the CL method

II.i The estimated outstanding claims

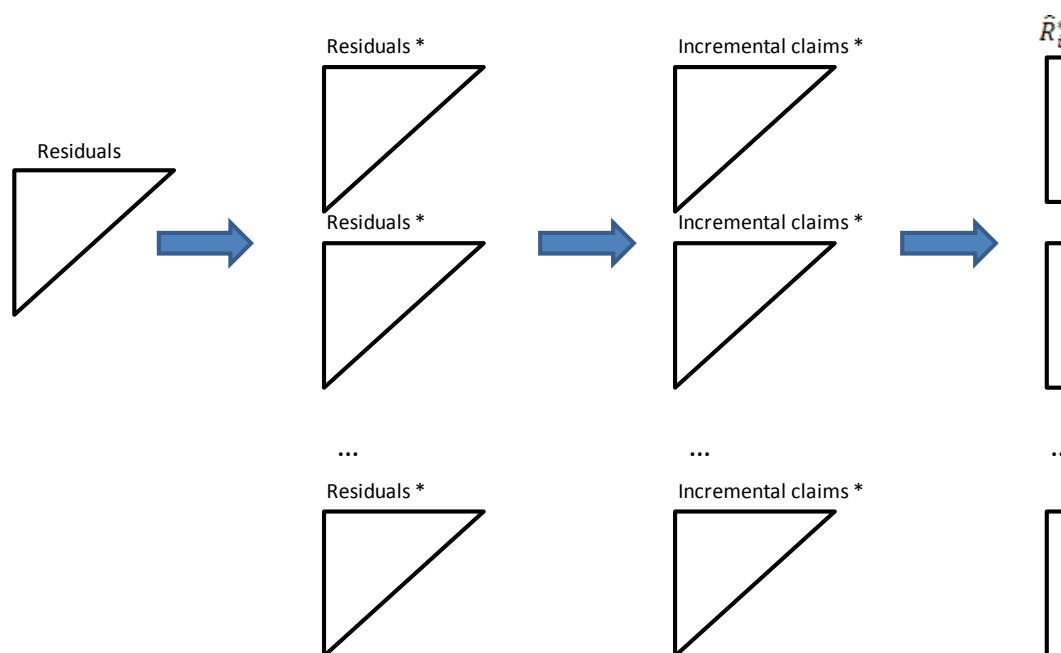
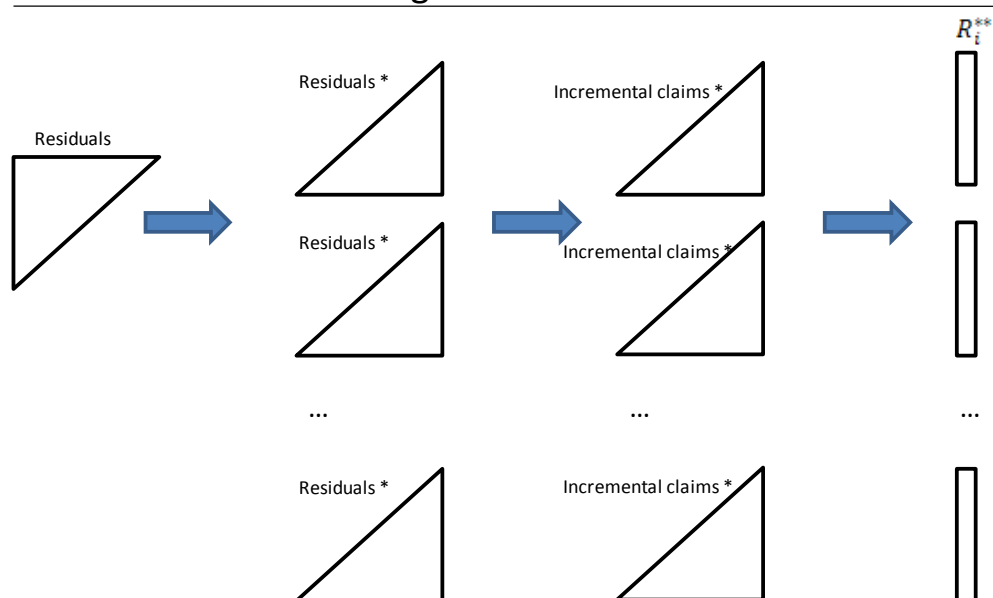


Figure 5

Bootstrap technique for the CL method

II.ii The true outstanding claims



$$pe_i^{**} = R_i^{**} - \hat{R}_i^*$$

Figure 6

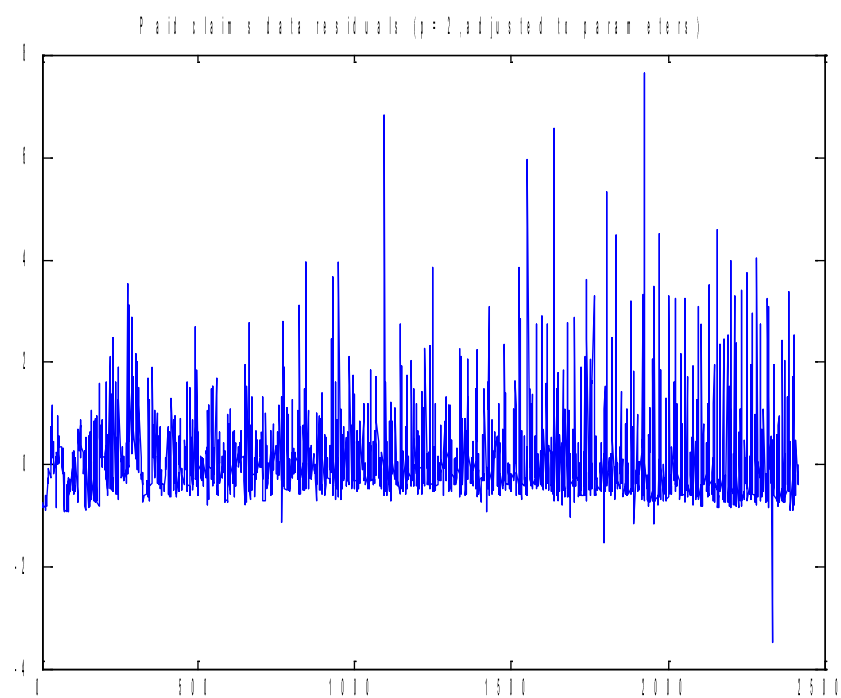


Figure 7. Paid claims data residuals ($p = 2$, adjusted to parameters).

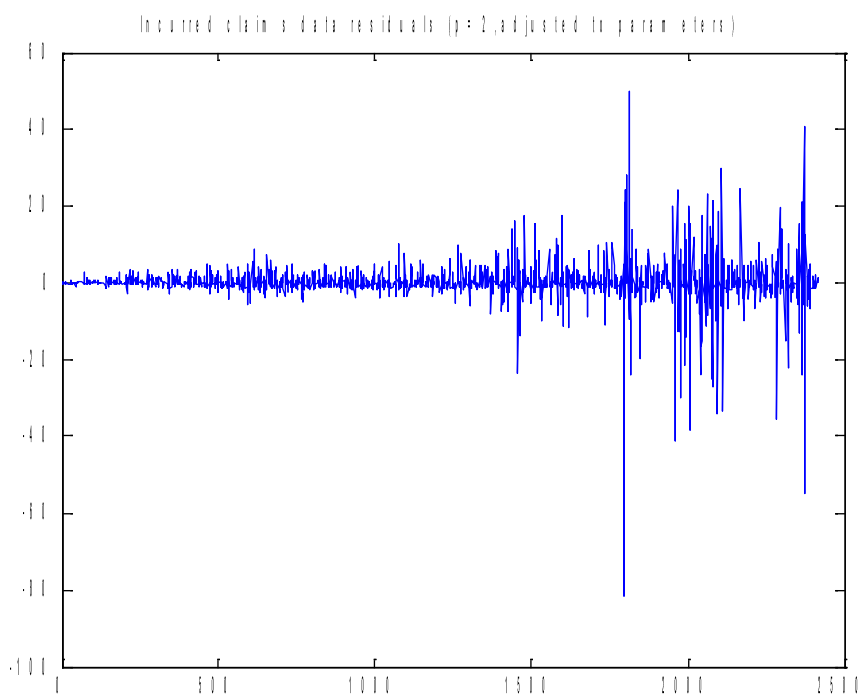


Figure 8. Incurred claims data residuals ($p = 2$, adjusted to parameters).

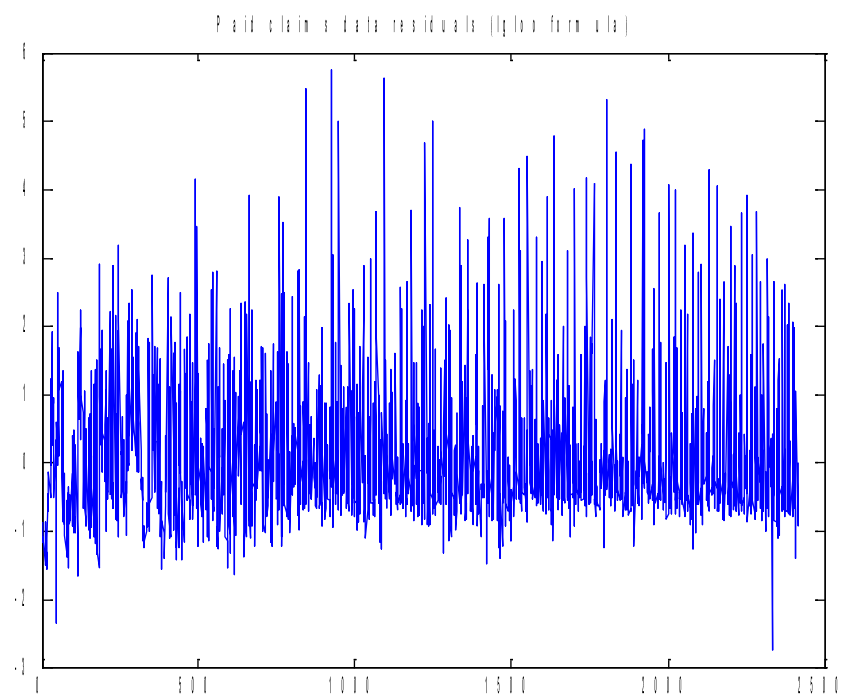


Figure 9. *Paid claims data residuals (lgloo formula).*

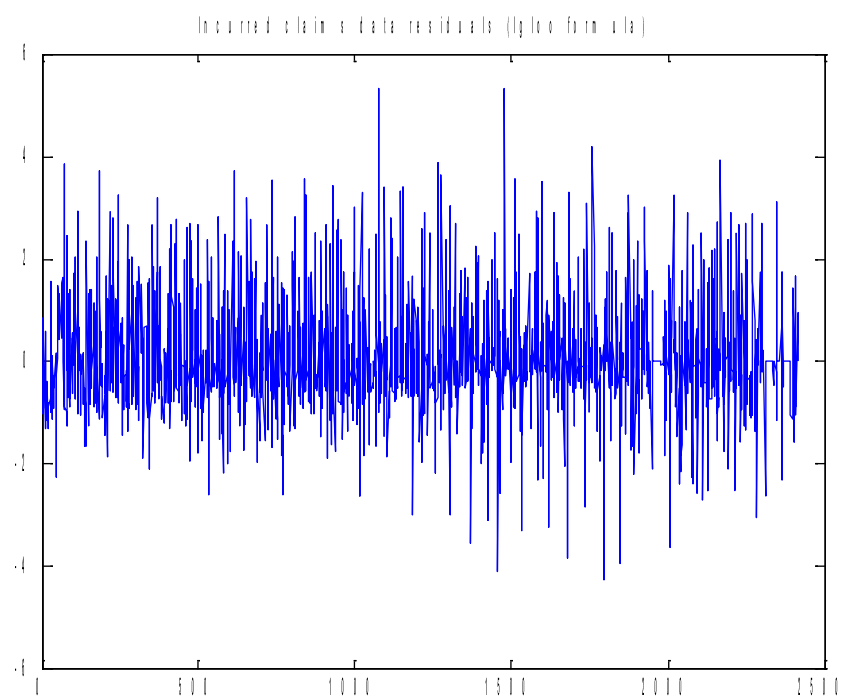


Figure 10. *Incurred claims data residuals (lgloo formula).*

3.2 The Bornhuetter-Ferguson Method

Using the notations introduced in the section 2, the BF method will be presented. Since it is important to see how this method deals with the weaknesses of the CL method, the presentation will start by showing the differences.

CL uses development factors $\hat{f}_k$ in order to project the current claims amount $C_{i,n+1-i}$ to the ultimate

$$\hat{U}_i^{CL} = \hat{C}_{i,n} = C_{i,n+1-i} \hat{f}_{n+2-i} \dots \hat{f}_{n-1}$$

Therefore, the CL reserve is

$$\hat{R}_i^{CL} = \hat{U}_i^{CL} - C_{i,n+1-i} = C_{i,n+1-i} (\hat{f}_{n+2-i} \dots \hat{f}_{n-1} - 1)$$

Notice that the reserve strongly depends on the current claims amount. It might very well happen so that the current claims amount is 0, then the reserve for that particular accident quarter will be estimated to be 0, which might lead to very wrong results.

This weakness of the CL method is avoided in the BF reserve estimate (Mack, 2008)

$$\hat{R}_i^{BF} = \hat{U}_i (1 - \hat{z}_{n+1-i}) \quad (7)$$

where $\hat{U}_i = v_i \hat{q}_i$ which is clearly independent from the current claims amount $C_{i,n+1-i}$. In the formula $\hat{U}_i$ is the prior estimate of ultimate claims, computed with a prior estimate $\hat{q}_i$ for the ultimate claims ratio $q_i = U_i/v_i$ of the accident period i , and $z_k \in [0,1]$ is the estimated percentage of the ultimate claims amount that is expected to be known after development period k . We see now that the BF reserve estimate is strongly dependent on the choice of prior ultimate, but not on the current claims amount.

To choose the proper prior ultimate loss ratios is the most important task in the BF method, but still not enough for calculating the BF reserve – the development pattern is also needed. Practitioners usually select the CL development pattern and construct z_k 's in the following way:

$$\hat{z}_n = \hat{f}_\infty^{-1}, \hat{z}_{n-1} = (\hat{f}_n \hat{f}_\infty)^{-1}, \dots, \hat{z}_1 = (\hat{f}_2 \dots \hat{f}_n \hat{f}_\infty)^{-1}$$

with $\hat{f}_\infty$ being the selected tail factor. This selection contradicts to the BF fundamental assumption that the reserve estimate does not depend on the current claims amount. Mack (2008) in his article

presents a different approach to estimate the development factor for the BF method, which makes the BF method a stand-alone reserving method. The latter approach was chosen to be used in this work.

The main problem of the thesis being the prediction error of the reserve estimate brings us to presenting the stochastic model for the BF method. The required prior belief in the ultimate claims strongly offers to use the Bayesian model. Verrall (2001) showed how Bayesian models within the framework of generalized linear model can be applied to claims reserving. Very recently Mack (2008) suggested a frequentist approach stochastic model, constructed in a similar manner as the one for the CL. The Mack (2008) model was chosen to be investigated in the thesis. In the following the Mack model and the prediction formulae will be presented.

3.2.1 Stochastic Model

First, it is assumed that all incremental claims $S_{i,k}$ of the same accident quarter are independent, as well as the accident quarters themselves. This assumption rises up much questioning, because in practice it is very rarely so. Nevertheless, the assumption must be satisfied in order to obtain the formula to calculate the prediction error. Therefore, we obey and assume the independence.

BF1 All increments $S_{i,k}$ are independent.

Since the reserve estimate is understood as the difference between the expected ultimate and the latest known claims amount, the second assumption follows from the BF reserve formula and claims that $E(C_{i,k}) = x_i z_k$. This can also be expressed for the incremental claims: $E(S_{i,k}) = x_i y_k$, $1 \leq i \leq n$ and $1 \leq k \leq n + 1$. Since the x_i and y_i parameters are unique only up to a constant factor, without loss of generality it is possible to assume $y_1 + y_2 + \dots + y_n + y_{n+1} = 1$. This yields to $E(U_i) = E(S_{i,1} + \dots + S_{i,n+1}) = E(x_i y_1 + \dots + x_i y_{n+1}) = x_i$ and therefore the parameter x_i can be considered as a measure of volume for accident period i .

BF2 There are unknown parameters x_i, y_k with $E(S_{i,j}) = x_i y_k$ and $y_1 + \dots + y_{n+1} = 1$.

Third, the variance of the ultimate claims $Var(U_i)$ should be proportional to x_i , or $Var(\frac{U_i}{x_i})$ proportional to $1/x_i$. This is the usual assumption for the influence of the volume on the variance. Since we assumed earlier that all the increments are independent, the variance of the incremental amount must also be proportional to x_i .

BF3 There are unknown proportionality constants s_k^2 with $Var(S_{i,j}) = x_i s_k^2$.

Then $Var(U_i) = Var(S_{i,1} + \dots + S_{i,n+1}) = Var(S_{i,1}) + \dots + Var(S_{i,n+1}) = x_i(s_1^2 + \dots + s_{n+1}^2)$ is proportional to x_i as intended.

To confirm that the assumptions work for the BF method we check:

$$E(R_i) = E(S_{i,n+2-i} + \dots + S_{i,n+1}) = x_i(y_{n+2-i} + \dots + y_{n+1}) = x_i(1 - z_{n+1-i})$$

with $z_k = y_1 + \dots + y_k$. The latter shows that the expected claims reserve has the same form as the BF reserve estimate (7). Since all the increments are independent,

$$\begin{aligned} Var(R_i) &= Var(S_{i,n+2-i} + \dots + S_{i,n+1}) = Var(S_{i,n+2-i}) + \dots + Var(S_{i,n+1}) \\ &= x_i(s_{n+2-i}^2 + \dots + s_{n+1}^2) \end{aligned}$$

Having the stochastic model for the BF method, the reserve estimate itself will be first calculated, and then formulae for the prediction error of the estimate will be presented.

3.2.2 Parameter Estimation

First step to a good reserve estimate is a proper choice of prior ultimate claims. The most important when making this choice is to bear in mind that the claims amount known so far should not be the main basis for the estimate $\hat{x}_i$. There are a few ways of obtaining the priors. In practice the estimation is mainly based on additional pricing and market information. Unfortunately, this information is not always available. Then it is not that uncommon among practitioners to use the CL ultimate claims as basis for the prior estimates. It is also possible, having premiums and run-off triangle data $\{v_i, C_{i,k}\}$, derive the prior loss ratios. The procedure of derivation is described in Mack (2006) and was implemented in the calculations for this work.

Since the raw claims data was available in the analysis of the personal accident claims, another approach was tried for computing the prior loss ratios. A distribution for the homogeneous classes of claims data (paid or incurred) was planned to be fitted, and the prior loss ratios would be simulated. Unfortunately, it appeared to be very difficult to fit a really good distribution for the data. The incremental claims triangles appeared to be rather sparse, and no distribution was found to be able to simulate enough zeroes. As a result, too large prior ultimate claims were simulated. Finally, it was decided not to develop this approach further due to the unrealistic reserve estimates it provided. Instead, to get the feeling of the prediction error when the BF method is used in practice, the prior ultimate claims were estimated as the CL ultimates, both raw and after the premiums' indexation.

Now, assuming that the prior loss ratios are estimated, denote it $\hat{U}_i$, the following task is to estimate the two unknown parameters, y_k and s_k^2 . In Mack (2008) two ways of estimating these parameters

are discussed – one mostly based on actuary's experience, and the other based on theoretic knowledge. In the numerical calculations the theoretical way was used, therefore, this method is presented here.

First, it is important to mention that due to the lack of data it is not possible to estimate the tail ratio y_{n+1} without further assumptions. One way to get it is from similar portfolios where the claims experience of later development is available. Another way is to extrapolate the series $\hat{y}_1, \dots, \hat{y}_n$. Similarly, an estimate of s_n^2 could be also obtained by extrapolation. Having this in mind, an iterative procedure was constructed in Mack (2008). We choose the following starting values:

$$\hat{y}_k = \frac{\sum_{i=1}^{n+1-i} S_{i,k}}{\sum_{i=1}^{n+1-i} \hat{U}_i}, 1 \leq k \leq n \quad (8)$$

The next step is to decide on the formula for a smoothing regression. Here we chose the Mack's (2008) offer $\ln(\hat{y}_k) = \alpha - \beta k$ for k above some $k_1 < n$. The smoothened regression is then extrapolated until some final development period $k_2 > n$. Now, the $\hat{s}_k^2$ can be calculated using the smoothened $\hat{y}_k$ for $k > k_1$:

$$\hat{s}_k^2 = \frac{1}{n-k} \sum_{i=1}^{n+1-k} (S_{i,k} - \hat{U}_i \hat{y}_k)^2 / \hat{U}_i \quad (9)$$

The resulting values should now be kept fixed and used in the following minimization of Q :

$$Q = \sum_{i=1}^n \sum_{k=1}^{n+1-i} \frac{(S_{i,k} - \hat{U}_i \hat{y}_k)^2}{\hat{U}_i \hat{s}_k^2} \quad (10)$$

under the two constraints:

$$\begin{cases} \hat{y}_1 + \dots + \hat{y}_n = 1 - \hat{y}_{n+1} \\ \hat{y}_1 + \dots + \hat{y}_{k_1} + \exp(\alpha - \beta(k_1 + 1)) + \dots + \exp(\alpha - \beta k_2) = 1 \end{cases}$$

After the minimization is done, the selection for all $\hat{y}_k^*$ is ready: the values for $k = 1, \dots, k_1$ are obtained directly, those for $k = k_1 + 1, \dots, n$ are taken from the smoothing regression and $\hat{y}_{n+1}$ is obtained by adding up the extrapolated values of the regression up to development period k_2 . Using the latest selection $\hat{y}_k^*$ we calculate $\hat{s}_k^{2'}$ s again with the (9) formula. Now $\ln(\hat{s}_k^{2'})$ for $k > k_1$ should be plotted against $|\hat{y}_k^*|$ or $\ln(|\hat{y}_k^*|)$ in order to select appropriate values for $\hat{s}_k^{2*}$, especially for $k = n$ and $k = n + 1$. The iteration could be continued, but Mack (2008) indicates that the continuation would not change values much at all; therefore, the latest selections are taken as estimates of the parameters.

The goal – the resulting development pattern being different from the CL development pattern – is achieved. Now, the BF reserve estimate can be calculated:

$$\hat{R}_i^{BF} = \hat{U}_i(\hat{y}_{n+2-i}^* + \dots + \hat{y}_{n+1}^*) = \hat{U}_i(1 - \hat{z}_{n+1-i}^*)$$

with $\hat{z}_k^* = \hat{y}_1^* + \dots + \hat{y}_k^*$. Mack (2008) lists 6 properties of the estimated parameters:

- $\hat{y}_1^*, \dots, \hat{y}_n^*, \hat{y}_{n+1}^*$ are pairwise (slightly) negatively correlated as they have to add up to unity.
- $\hat{y}_1^*, \dots, \hat{y}_n^*, \hat{y}_{n+1}^*$ and therefore also $\hat{z}_1^*, \dots, \hat{z}_n^*$ are practically independent from $\hat{U}_1, \dots, \hat{U}_n$ as the latter do not really influence the size of any $\hat{y}_k^*$ because these have to add up to unity in any case and because of selections and regressions used.
- $\hat{R}_i^{BF}$ and R_i are independent (due to (BF1)).
- $E(\hat{U}_i) = E(U_i) = x_i, 1 \leq i \leq n$.
- $E(\hat{y}_k^*) = y_k, 1 \leq k \leq n+1$, and therefore $E(\hat{z}_k^*) = z_k, 1 \leq k \leq n+1$.
- $E(\hat{s}_k^{2*}) = s_k^2, 1 \leq k \leq n+1$.

From the unbiasedness of the unknown parameters it follows:

$$E(\hat{R}_i^{BF}) = E(\hat{U}_i)E(1 - \hat{z}_{n+1-i}^*) = x_i(1 - z_{n+1-i}) = E(R_i)$$

which means that the BF reserve estimate is unbiased.

3.2.3 Calculating the Prediction Error

Having stochastic model assumptions for the BF method stated and the procedures of estimating unknown parameters derived, the formula for computing the prediction error of the reserve estimate for each accident quarter as well as total will be given. All formulae presented in this section were derived in Mack (2008). In Appendix II the proofs (Mack, 2008) with minor clarifications added are presented.

As discussed earlier, we are interested in the mean square error of the prediction given the data observed so far, i.e. we are interested only in future variability. Therefore, the mean square error of the prediction of any reserve estimate $\hat{R}_i$ is defined by

$$mse(\hat{R}_i) = E\left((\hat{R}_i - R_i)^2 \middle| S_{i,1}, \dots, S_{i,n+1-i}\right)$$

Since $R_i = S_{n+2-i} + \dots + S_{i,n+1}$ is independent from $S_{i,1}, \dots, S_{i,n+1-i}$ according to (BF1), also, the BF reserve estimate can be taken as independent of the increments. Therefore,

$$mse(\hat{R}_i^{BF}) = E\left((\hat{R}_i^{BF} - R_i)^2\right) = Var(\hat{R}_i^{BF} - R_i) + (E(\hat{R}_i^{BF}) - E(R_i))^2 = Var(\hat{R}_i^{BF}) + Var(R_i)$$

Notice, the mean square error of prediction of the BF reserve estimate is the sum of the squared estimation error $Var(\hat{R}_i^{BF})$ and of the squared process error $Var(R_i)$.

Theorem 2. Under the assumptions (BF1), (BF2), and (BF3) the mean square error $mse(\hat{R}_i)$ can be estimated by

$$mse(\hat{R}_i^{BF}) = \hat{U}_i(\hat{s}_{n+2-i}^2 + \dots + \hat{s}_{n+1}^2) + \left(\hat{U}_i^2 + (s.e.(\hat{U}_i))^2 \right) (s.e.(\hat{z}_{n+1-i}^*))^2 + (s.e.(\hat{U}_i))^2 (1 - \hat{z}_{n+1-i}^*)^2 \quad (11)$$

The estimation error can be examined more thoroughly.

$$(s.e.(\hat{R}_i^{BF}))^2 = \left(\hat{U}_i^2 + (s.e.(\hat{U}_i))^2 \right) (s.e.(\hat{z}_{n+1-i}^*))^2 + (s.e.(\hat{U}_i))^2 (1 - \hat{z}_{n+1-i}^*)^2$$

Dividing it by the prior ultimate claims amount the following can be noticed:

$$s.e.(\hat{R}_i^{BF})/\hat{U}_i \approx s.e.(\hat{z}_{n+1-i}^*) \text{ for } \hat{z}_{n+1-i}^* \text{ close to } 1$$

$$s.e.(\hat{R}_i^{BF})/\hat{U}_i \approx s.e.(\hat{U}_i)/\hat{U}_i \text{ for } \hat{z}_{n+1-i}^* \text{ close to } 0$$

This can be interpreted as follows: for the accident quarters where there is almost no development left the uncertainty of the initial ultimate claims estimates is directly transferred to the reserve estimate.

By now we have presented how to compute the error of prediction for the reserve estimates for each accident period. Just like in the CL method, it might be so, that one is interested in the overall reserve estimate and its variability. The overall BF reserve is just a sum of all reserve estimates for each accident period $\hat{R}^{BF} = \hat{R}_1^{BF} + \dots + \hat{R}_n^{BF}$. Again, when summing the prediction error the covariance among the reserve estimates must be taken into the account.

Corollary 2. With the assumptions and notations of Theorem 2 the mean squared error of the overall reserve estimate $\hat{R}^{BF} = \hat{R}_1^{BF} + \dots + \hat{R}_n^{BF}$ can be estimated by

$$mse(\hat{R}^{BF}) = \sum_{i=1}^n \hat{U}_i(\hat{s}_{n+2-i}^2 + \dots + \hat{s}_{n+1}^2) + \sum_{i=1}^n Var(\hat{R}_i^{BF}) + 2 \sum_{i < j} Cov(\hat{R}_i^{BF}, \hat{R}_j^{BF}) \quad (12)$$

3.2.4 Bootstrapping the Estimation Error

The bootstrap methodology for the BF method has not yet been introduced in any literature that was available during the time of writing this work. Some underlying thoughts, though, were found in

England & Verrall (2004) presentation. Adding this, own understanding of the method and the bootstrap methodology for the CL method, the procedure presented below was derived. Unfortunately, the furthest we were able to arrive is the estimation error, not the process error. The problem here is that all the derived procedures for the CL method bootstrapping computes the process error by building B lower triangles. The BF method has no way of building a lower triangle. It was decided to stop at this point, and leave the estimation of the process error for the BF method outside the scope of the thesis. The constructed procedure for bootstrapping the estimation error of the BF method is the following:

- I. The preliminaries
 - Estimate the development pattern $\hat{z}_1, \hat{z}_2, \dots, \hat{z}_{n-1}$ for the triangle of the data
 - Calculate the fitted values $\hat{m}_{ij}, i = 1, 2, \dots, n, j = 1, 2, \dots, n + 1 - i$
 - Calculate the residuals
 - Estimate the prior ultimate claims $\hat{U}_1, \hat{U}_2, \dots, \hat{U}_n$ for the triangle of the data
 - Calculate the outstanding claims $\hat{R}_i = \hat{U}_i(1 - \hat{z}_{n-i+1})$ and $\hat{R} = \hat{R}_2 + \dots + \hat{R}_n$
- II. Bootstrap world (to be repeated B times)
 - i. The estimated outstanding claims
 - Resample the residuals obtained in the first stage using replacement
 - Create pseudo-data by solving the residuals' formula backwards with $\hat{m}_{ij}, i = 1, 2, \dots, n, j = 1, 2, \dots, n + 1 - i$
 - Estimate the development pattern and the prior ultimate claims with the pseudo-data
 - Calculate the estimated outstanding claims $\hat{R}_i^* = \hat{U}_i^*(1 - \hat{z}_{n-i+1}^*)$ and $\hat{R}^* = \hat{R}_2^* + \dots + \hat{R}_n^*$

To be able to fit the values the development factors from the estimated development pattern have to be calculated using the formula backwards:

$$\hat{z}_n = \hat{f}_\infty^{-1}, \hat{z}_{n-1} = (\hat{f}_n \hat{f}_\infty)^{-1}, \dots, \hat{z}_1 = (\hat{f}_2 \dots \hat{f}_n \hat{f}_\infty)^{-1}$$

Since the minimization of Q (10) is changing the development pattern estimates extremely little, we have decided not to implement minimization in every bootstrap loop to save the time of computations. The residuals were computed following the formula Trygg-Hansa's actuaries use in the EMB Igloo Professional. Also, the prior ultimate claims were estimated in three ways: as Mack (2006), as the CL ultimate claims, and as the CL ultimate claims after the premiums' indexation. Moreover, two approaches of bootstrapping will be implemented. The first one – the prior ultimates are estimated from the original data triangle and then used as a constant for each bootstrap loop. The

other approach, when the prior ultimates are estimated in each bootstrap loop separately from a pseudo-triangle data.

4. Data Analysis

4.1 Product Presentation

The data to be analyzed is personal accident insurance for children sold at Trygg-Hansa. One part of the data comprises the run-off product, i.e. not anymore on the market, and the other part is product which is being sold at the moment. Depicting shortly what the product gives to the customer is the following:

- Illness and accident insurance that is valid 24 hours a day.
- Insurance also covers invalidity, costs for hospital, and life.
- Quick cover for certain diseases, ex. Cancer, etc.
- Possibility to sign for different sizes of insurance amount.
- Valid up to and including 25 years old age. Then it is transferred to a life, illness, accident insurance for an adult without additional health check.
- Premium is adjusted according to the child's age.
- The new product covers both medical and economic disability.
- The old product covers either economic or medical disability (which is the best for the customer).
- Monthly compensation to parents and to the child after 18 years of age.
- Compensation for parents care in the event of long-term sickness and injury until the child is 30 year old (25 in the old product) or economic disability can be determined.

In the above, medical disability covers compensation in reduced functional ability, and economic disability gives compensation for reduced work ability.

The claims data presented in the thesis were multiplied by a factor to keep the confidentiality of the company.

4.2 Results

Having the claims data summarized into a run-off triangle, the reserving analysis can be implemented. In the following two types of claims data will be analyzed, paid and incurred claims. Paid claims represent the money the insurance company has paid to its clients, and incurred claims represent the money which was estimated to be claimed from the company. Reserving both paid and incurred gives a deeper insight to an actuary about the future. Usually, the resulting ultimate claims are not the same for paid and incurred data. A subjective choice should be made on which to take as the concluding answer. Sometimes the two estimates can be combined and a weighted average would be taken. In Trygg-Hansa, as a rule, the incurred claims data is the main basis for the concluding results.

The *Results* section is divided into three parts: the CL method, the BF method, and the comparison study. For each method theoretical and bootstrap calculations were implemented. First, the created procedures were checked with the standard data triangle from Taylor & Ashe (1983). After the procedures were confirmed to be correct the PA Children data was analyzed.

The bootstrap procedures outlined in section 3 were implemented with $B = 10\,000$. The data for bootstrapping was taken starting with the accident quarter 1991 Q1 in order to have a full triangle. As is indicated in Björkwall (2009), the estimated outstanding claims from the bootstrap world can be seen as the estimation error component; the true outstanding claims from the bootstrap world can be seen as the process error components; and finally, the predictive distribution of the reserve estimate is the reserve estimate plus the prediction error (Björkwall, 2009). As for the BF method, only the estimation error was obtained, which will be compared to the theoretical estimation error of the BF method, and also to the bootstrapped CL estimation error. In the presentation of results you will find *95 percentile* columns. This is the 95th percentile of the predictive reserve distribution.

Theoretical and bootstrap PA Children data results will be presented for both methods.

4.2.1 The Chain-Ladder Results

4.2.1.1 Stochastic Assumptions

As was explained in the section 3.1.2, first the assumptions for the data have to be checked. The linearity between incurred data in two adjacent development periods is clearly seen from the graphs in Appendix III.a. Notice that even when there are not many points of data (late development quarters), the points are distributed clearly on a line. The assumption of linearity (CL1) is strongly confirmed for the incurred data. The graphs for the paid cumulative claims data can be found in Appendix III.b. The paid data also confirms linearity assumption (CL1).

Secondly, the assumption (CL3) about the variance of data has to be checked. The graphs with the residuals plotted against the cumulative data are presented in the Appendix III.c for the incurred claims data and in the Appendix III.d for the paid claims data. In order to accept the assumption, the graphs should not indicate any trends. For both, incurred and paid claims data, the residuals for the first development periods have a tendency to be more negative for the older accident quarters and positive for the more recent accident quarters. This tendency disappears after approximately 20 development periods. For the latest development periods due to so little data one can assume that the residuals are randomly distributed. Since in the most of the development periods calculated residuals seem to have no trends, we assume that the (CL3) assumption is accepted.

Finally, the independence assumption has to be verified. For that the test derived in Mack (1994) was implemented. The resulting test statistics with corresponding 95% confidence interval bounds are presented in the Table 1 below. We see that both, Z_{inc} and Z_{paid} are inside their confidence intervals, therefore the test of accident periods' independence gave a positive result.

	Z	E(Z)	Var(Z)	Lower bound	Upper bound
Incurred claims	1870	2203,1	441,6	1319,9	3086,3
Paid claims	1503	2198,5	441,58	1315,3	3081,6

Table 1. The (CL2) assumption test.

4.2.1.2 Theoretical Prediction Error

Now that all the assumptions of stochastic model for the CL method are checked and correct, the actual model can be applied, and the prediction error of the result can be calculated. This was done for both, incurred and paid claims data. In the Tables 2 and 3 below you can see the CL result on the data for 16 recent accident quarters, as well as the total result. Also, in the Figure 11 the prediction error for accident quarters since 1986Q1 up to the end are presented graphically for both, incurred claims and paid claims data. Comparing the results on incurred and paid claims data a clear

conclusion that the CL method gives much more exact results on paid data can be drawn, especially for the earlier accident quarters. Also, important to notice that the resulting ultimate claims from the CL on the paid data are higher than the ones on the incurred data, which, together with a smaller variability means, that it is more likely to get higher ultimate claims in the future than optimistic, lower ones. The prediction error is decreasing over the accident periods, but for the most recent accident quarters the error increases due to very little data observed.

Acc. Qrt	Premiums	Ultimate	Case Reserve	IBNR Reserve	LR	s.e. (R)	Theoretical s.e. %	ResQ s.e. %
2005 Q1	17 901 523	13 411 927	4 033 476	4 181 564	74,92%	1 871 732	44,76%	37,21%
2005 Q2	18 028 323	14 351 004	4 128 687	4 631 488	79,60%	2 012 477	43,45%	35,28%
2005 Q3	18 837 130	15 061 218	3 949 018	5 038 087	79,95%	2 122 975	42,14%	33,79%
2005 Q4	19 658 343	12 590 843	3 380 709	4 372 032	64,05%	1 784 862	40,82%	35,72%
2006 Q1	32 484 747	17 948 519	5 024 257	6 480 849	55,25%	2 556 357	39,44%	29,87%
2006 Q2	21 087 053	16 151 979	4 855 576	6 076 791	76,60%	2 315 395	38,10%	30,61%
2006 Q3	22 340 739	15 251 531	4 286 522	5 992 760	68,27%	2 198 846	36,69%	30,61%
2006 Q4	23 415 037	13 453 508	3 719 061	5 536 366	57,46%	1 961 322	35,43%	31,81%
2007 Q1	24 149 928	14 690 463	3 954 223	6 352 769	60,83%	2 167 457	34,12%	30,05%
2007 Q2	23 944 783	19 465 604	5 536 118	8 882 550	81,29%	2 893 966	32,58%	25,78%
2007 Q3	24 583 795	21 389 175	6 487 846	10 353 575	87,01%	3 215 206	31,05%	24,24%
2007 Q4	25 252 561	21 098 094	6 876 834	10 908 403	83,55%	3 214 545	29,47%	23,83%
2008 Q1	25 696 884	22 179 253	6 625 982	12 365 242	86,31%	3 472 161	28,08%	23,39%
2008 Q2	25 794 583	23 281 603	6 785 309	14 118 786	90,26%	3 958 340	28,04%	25,91%
2008 Q3	26 215 594	28 410 351	7 047 765	19 256 555	108,37%	5 327 942	27,67%	26,67%
2008 Q4	26 960 369	28 573 686	4 253 631	23 625 938	105,98%	8 013 009	33,92%	41,34%
Total	887 313 233	751 914 300	192 279 083	227 728 625	84,74%	26 248 977	11,53%	13,19%

Table 2. *The Chain-Ladder result on incurred claims data.*

Acc. Qrt	Premiums	Ultimate	Case Reserve	IBNR Reserve	LR	s.e. (R)	Theoretical s.e. %	ResQ s.e. %
2005 Q1	17 901 523	14 757 719	4 033 476	5 527 356	82,44%	1 770 739	18,52%	14,84%
2005 Q2	18 028 323	16 883 129	4 128 687	7 163 612	93,65%	2 056 032	18,21%	14,23%
2005 Q3	18 837 130	19 630 876	3 949 018	9 607 745	104,21%	2 414 529	17,81%	13,40%
2005 Q4	19 658 343	16 859 463	3 380 709	8 640 652	85,76%	2 102 898	17,49%	14,34%
2006 Q1	32 484 747	24 422 112	5 024 257	12 954 442	75,18%	3 073 880	17,10%	12,31%
2006 Q2	21 087 053	21 740 935	4 855 576	11 665 747	103,10%	2 769 424	16,76%	13,00%
2006 Q3	22 340 739	22 937 733	4 286 522	13 678 962	102,67%	2 971 431	16,54%	12,99%
2006 Q4	23 415 037	21 973 553	3 719 061	14 056 411	93,84%	2 898 601	16,31%	13,46%
2007 Q1	24 149 928	25 589 769	3 954 223	17 252 076	105,96%	3 411 587	16,09%	12,72%
2007 Q2	23 944 783	32 624 300	5 536 118	22 041 246	136,25%	4 399 123	15,95%	11,74%
2007 Q3	24 583 795	33 678 296	6 487 846	22 642 696	136,99%	4 794 837	16,46%	13,43%
2007 Q4	25 252 561	29 973 305	6 876 834	19 783 615	118,69%	4 506 927	16,90%	15,62%
2008 Q1	25 696 884	37 538 009	6 625 982	27 723 997	146,08%	7 277 345	21,19%	22,99%
2008 Q2	25 794 583	33 864 228	6 785 309	24 701 411	131,28%	7 340 621	23,31%	27,76%
2008 Q3	26 215 594	42 276 500	7 047 765	33 122 704	161,26%	9 996 409	24,88%	28,12%
2008 Q4	26 960 369	39 288 477	4 253 631	34 340 728	145,73%	11 335 538	29,37%	36,63%
Total	887 313 233	844 453 385	192 279 083	320 267 710	95,17%	31 381 917	6,12%	7,68%

Table 3. *The Chain-Ladder result on paid claims data.*

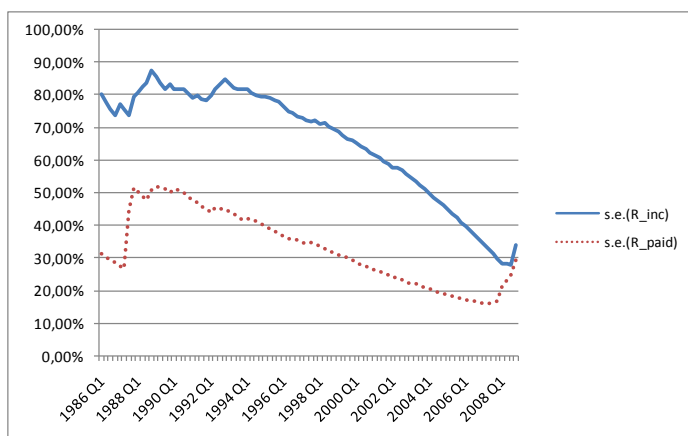


Figure 11. *The comparison of the prediction error between paid and incurred data.*

One can also compare the resulting prediction error from this calculation with the results given in EMB ResQ Professional. The reserve estimates are identical since the same CL method was used, but the prediction error given in the ResQ program is little different than the one computed here. This is due to some extra smoothing that the ResQ program does, which is not included in Mack (1993), the underlying basis for this work. For the comparison the standard error (s.e.) calculated in the ResQ program on the incurred claims data for the 16 recent quarters is presented in the last column of the Table 2. The same for the paid claims data is presented in Table 3. The comparison graphically can be seen in Figure 12 for the incurred claims data and Figure 13 for the paid claims data.

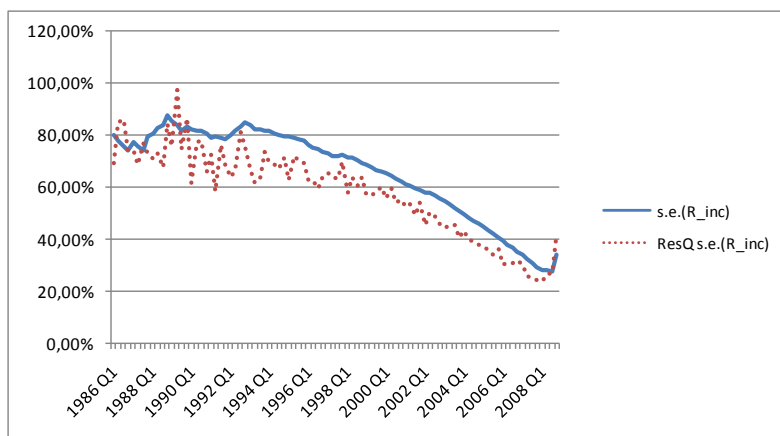


Figure 12. *The comparison of the prediction error between the thesis calculation and the ResQ result, incurred claims.*

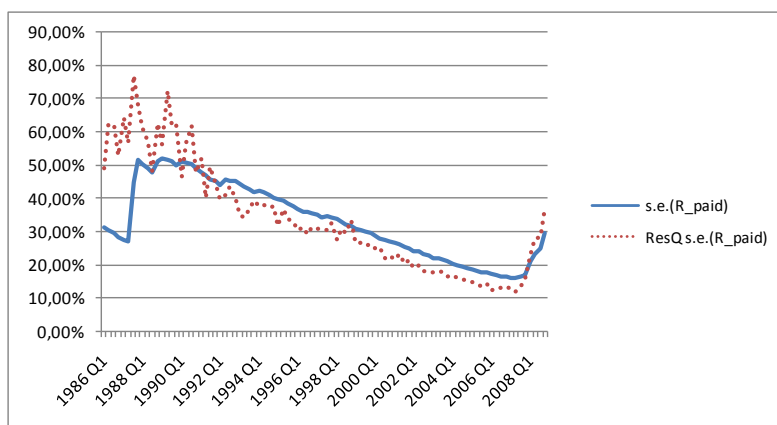


Figure 13. The comparison of the prediction error between the thesis calculation and the ResQ result, paid claims.

4.2.1.3 Bootstrapped Prediction Error

In the Table 4 the CL incurred claims reserve estimate with the 95th percentile of its predictive distribution and the prediction error in numbers and in percents are presented. Also, the theoretical prediction error is shown next to the bootstrapped one for comparison. The table includes only the most recent accident quarters' data. Whole time series of the bootstrapped prediction error together with the theoretical one is presented in the Figure 14. The same just for the paid claims data is given in the Table 5 and Figure 15.

The bootstrapped prediction error is smaller almost all the time for both, incurred and paid claims data, when calculated for each accident quarter separately. If we compare the total reserve prediction error, for paid claims data the bootstrapped prediction error is smaller than the theoretical one, but for the incurred claims the bootstrapped error is the higher one. A possible explanation to this is that the bootstrapping technique allows for stronger covariance among the incurred claims reserve estimates, but not among the paid claims reserve estimates.

The outstanding claims, estimated (a) and true (b), from the bootstrap world are given in the histograms in the Figure 16 for the incurred claims and in the Figure 17 for the paid claims. The histogram of the predictive distribution of the reserve estimate for both data is also presented (c). In the histograms one can see that the process error (b) is less spread out than the estimation error (a). The predictive reserve distribution (c) is as spread out as the estimation error (a). This holds for both, incurred claims and paid claims data. Also, from the histograms we can see that the prediction error of the paid claims reserve estimate is smaller than the one of the incurred claims reserve estimate.

Acc. Qrt	IBNR Reserve (000's)	95 percentile (000's)	s.e. (R) (000's)	s.e. %	Theoretical s.e. %
2005 Q1	4 182	5 995	1 777	42,50%	44,76%
2005 Q2	4 631	6 506	1 892	40,86%	43,45%
2005 Q3	5 038	6 936	2 031	40,31%	42,14%
2005 Q4	4 372	6 317	1 775	40,59%	40,82%
2006 Q1	6 481	8 470	2 496	38,52%	39,44%
2006 Q2	6 077	8 084	2 284	37,59%	38,10%
2006 Q3	5 993	8 008	2 185	36,46%	36,69%
2006 Q4	5 536	7 718	2 008	36,28%	35,43%
2007 Q1	6 353	8 556	2 226	35,04%	34,12%
2007 Q2	8 883	11 079	2 869	32,30%	32,58%
2007 Q3	10 354	12 697	3 175	30,66%	31,05%
2007 Q4	10 908	13 380	3 205	29,38%	29,47%
2008 Q1	12 365	15 161	3 564	28,83%	28,08%
2008 Q2	14 119	17 201	4 120	29,18%	28,04%
2008 Q3	19 257	23 223	5 419	28,14%	27,67%
2008 Q4	23 626	31 023	7 946	33,63%	33,92%
Total	227 729	240 366	44 045	19,58%	11,53%

Table 4. The bootstrapped prediction error and the 95th percentile of the predictive reserve distribution for the incurred claims data.

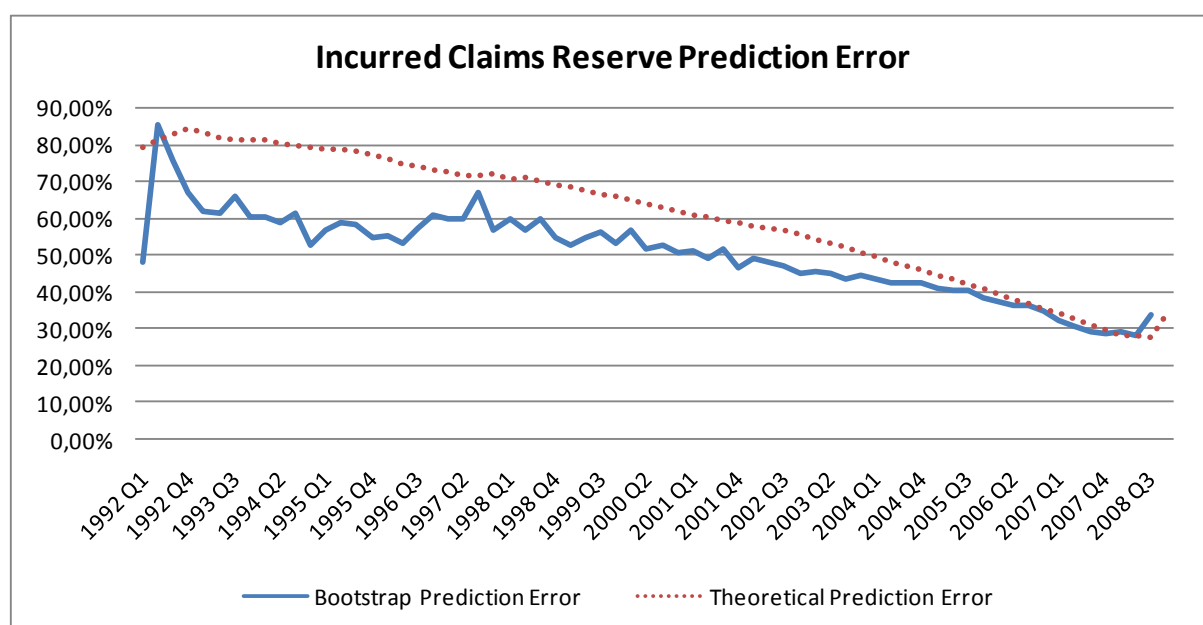


Figure 14. Incurred claims IBNR reserve prediction error.

Acc. Qrt	IBNR Reserve (000's)	95 percentile (000's)	s.e. (R) (000's)	s.e. %	Theoretical s.e. %
2005 Q1	5 527	11 556	1 316	13,77%	18,52%
2005 Q2	7 164	13 461	1 460	12,93%	18,21%
2005 Q3	9 608	15 953	1 638	12,08%	17,81%
2005 Q4	8 641	14 292	1 550	12,89%	17,49%
2006 Q1	12 954	20 875	1 984	11,04%	17,10%
2006 Q2	11 666	19 355	1 984	12,01%	16,76%
2006 Q3	13 679	21 095	2 178	12,12%	16,54%
2006 Q4	14 056	21 101	2 356	13,25%	16,31%
2007 Q1	17 252	25 134	2 818	13,29%	16,09%
2007 Q2	22 041	32 458	3 523	12,77%	15,95%
2007 Q3	22 643	34 571	3 932	13,50%	16,46%
2007 Q4	19 784	32 389	4 293	16,10%	16,90%
2008 Q1	27 724	40 635	4 695	13,67%	21,19%
2008 Q2	24 701	37 832	5 033	15,98%	23,31%
2008 Q3	33 123	48 067	6 724	16,74%	24,88%
2008 Q4	34 341	47 319	7 942	20,58%	29,37%
Total	320 268	363 611	30 121	5,93%	6,12%

Table 5. The bootstrapped prediction error and the 95th percentile of the predictive reserve distribution for the paid claims data.

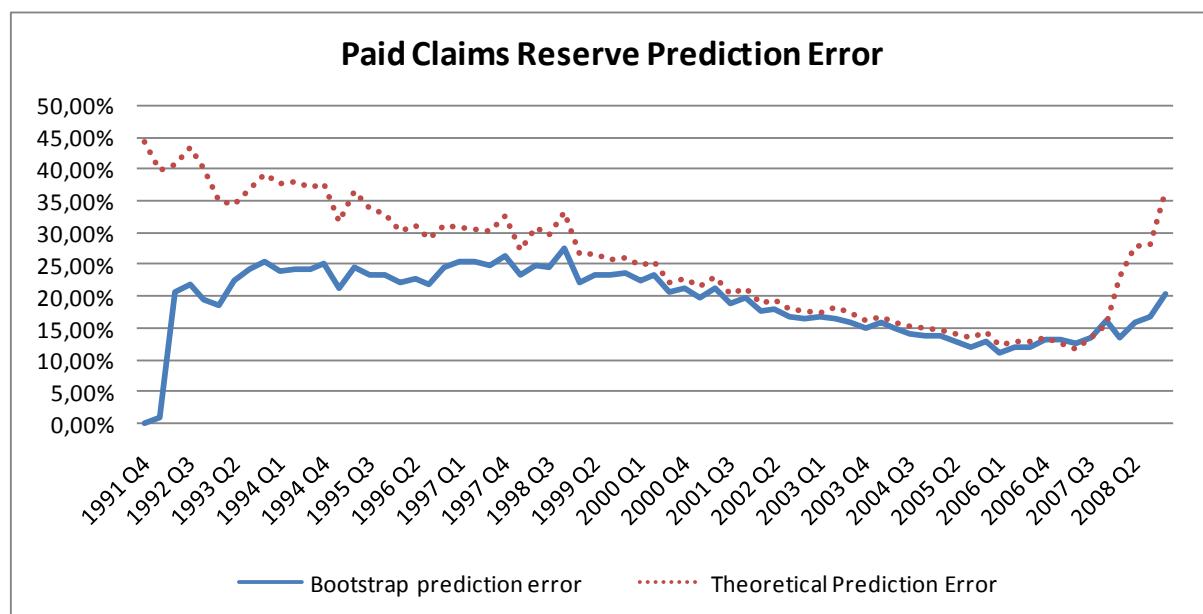


Figure 15. Paid claims IBNR reserve prediction error.

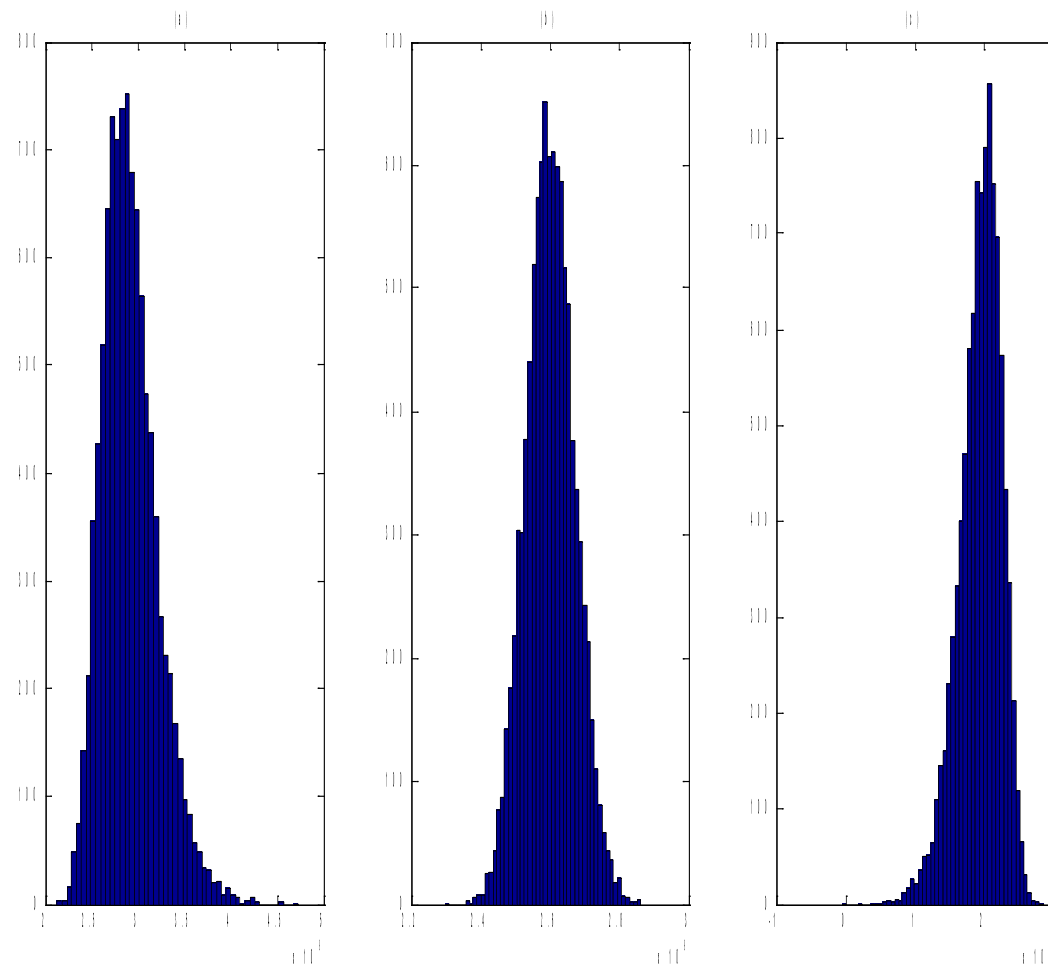


Figure 16. Density graphs of $\hat{R}^*$ (a), R^{**} (b), and $\tilde{R}^{**}$ (c) for the incurred claims data.

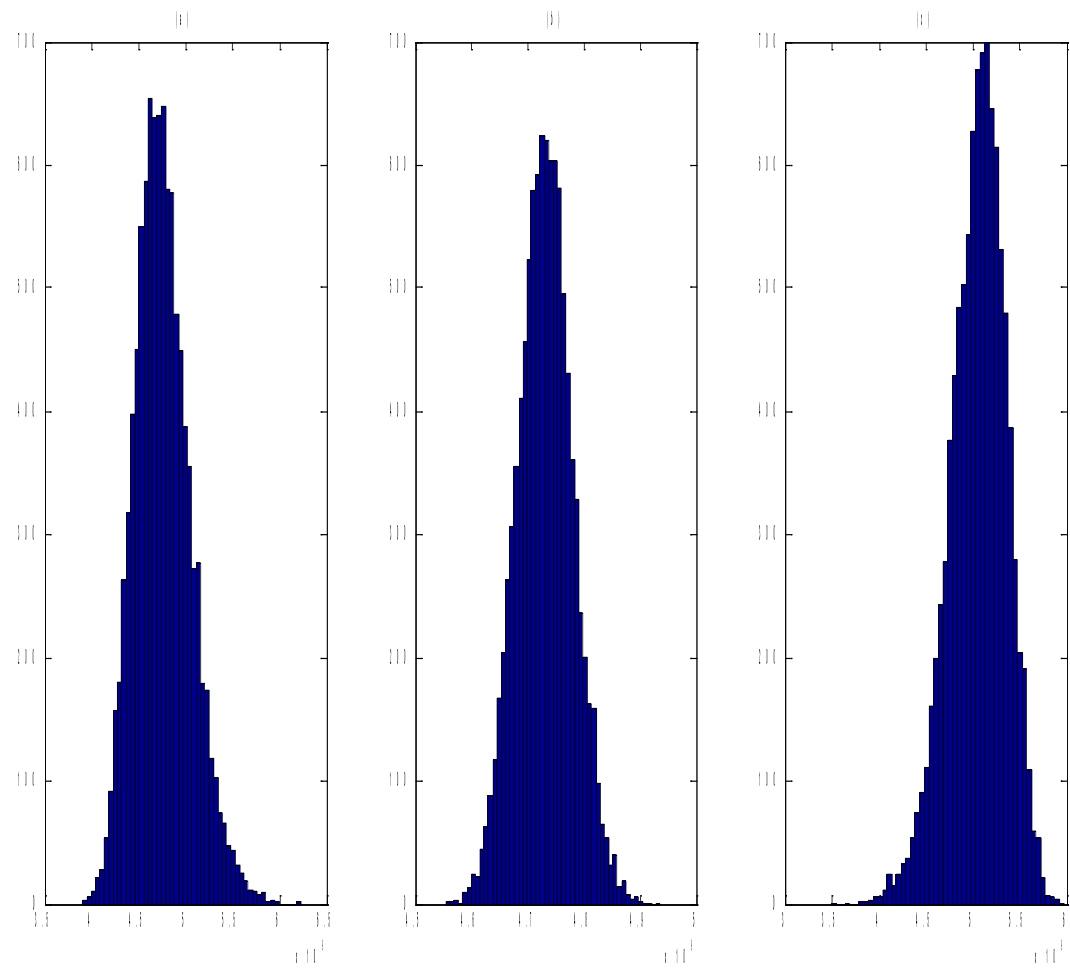


Figure 17. Density graphs of $\hat{R}^*$ (a), R^{**} (b), and $\tilde{R}^{**}$ (c) for the paid claims data.

4.2.2 The Bornhuetter-Ferguson Result

4.2.2.1 Estimation of Prior Ultimate Claims

Three ways of estimating the prior ultimate claims were chosen. The first way is a theoretical estimation according to the article by Mack (2006). The other two ways are estimations commonly used in practice. One of them is the CL method ultimate claims. The other one – the CL method ultimate claims with applied indexation of the premiums. The indexation brings the ultimate claims to a level which is more relative today. The two practical estimations are actually not following the main idea of the BF method, that the prior estimates are independent of the historical claims information. Nevertheless, these estimations were chosen because practically they are often used, and it is of interest what prediction error it has.

The theoretical estimation of prior ultimate claims includes the following steps:

1. Raw incremental loss ratio (ILR) at development quarter k : $ILR_k = \sum_{i=1}^{n+1-k} S_{i,k} / \sum_{i=1}^{n+1-k} v_i$
2. Raw on-level premium factor for accident quarter i : $r_i = C_{i,n+1-i} / (\frac{v_i}{\sum_{k=1}^{n+1-k} ILR_k})$
3. Selected on-level premium factor for accident quarter i (same for incurred and paid):

$$r_i^* = \sqrt{r_i^{paid} * r_i^{incurred}}$$
4. Adjusted average ILR at development quarter k : $ILR_k^{adj} = \sum_{i=1}^{n+1-k} S_{i,k} / \sum_{i=1}^{n+1-k} v_i r_i^*$
5. Selected average ILR at development quarter k : smoothed version of $ILR_k^{adj} = ILR_k^{adj*}$
6. *A priori* ULR for accident quarter i , with possibility to include tail ratio: $ULR_i = r_i^* * \sum_{k=1}^{n+1} ILR_k^{adj*}$
7. *A priori* estimate of ultimate losses for accident quarter i : $UL_i = ULR_i * v_i$

The result for each step of the theoretical estimation for most recent accident quarters is presented in the Table 6 below. You also find graphical presentation of the ILR, adjusted and selected, as well as the on-level premium factors, since 1986 Q1 (see Figure 18 (a)-(c)). The rate for the premiums has extreme values for accident quarters 2003 Q1 – Q3, because of very low premium data available. For the second half of the development periods the ILR has rather strong oscillations, which is smoothened by fitting the first order polynomial to the original adjusted ILR. The sum of selected ILR is seen as the ultimate loss ratio, which is later multiplied by the on-level premium factor to obtain ultimate loss ratios for each accident period.

Acc. Qrt	Premiums (000's)	Latest incurred (000's)	ILR (Incurred Claims)	Incurred rate	Latest paid (000's)	ILR (Paid Claims)	Paid rate	Chosen rate	Premiums *	Adjusted ILR (Incurred Claims)	Selected ILR (Incurred Claims)	Prior Ultimate Incurred Claims (000's)	Adjusted ILR (Paid Claims)	Selected ILR (Paid Claims)	Prior Ultimate Paid Claims (000's)
2005 Q1	17 902	9 230	0,01270	0,93360	5 197	0,01370	0,98550	0,95920	17 171	0,01260	0,01260	15 142	0,01360	0,01360	17 011
2005 Q2	18 028	9 720	0,01140	0,99900	5 591	0,01400	1,10400	1,05020	18 933	0,01140	0,01140	16 697	0,01390	0,01390	18 757
2005 Q3	18 837	10 023	0,01130	1,00730	6 074	0,01360	1,20790	1,10310	20 779	0,01120	0,01120	18 324	0,01340	0,01340	20 586
2005 Q4	19 658	8 219	0,01140	0,80890	4 838	0,01470	0,97130	0,88640	17 425	0,01140	0,01140	15 366	0,01460	0,01460	17 263
2006 Q1	32 485	11 468	0,01080	0,69840	6 443	0,01860	0,83090	0,76180	24 747	0,01090	0,01090	21 823	0,01880	0,01880	24 517
2006 Q2	21 087	10 075	0,00830	0,96600	5 220	0,02060	1,12460	1,04230	21 979	0,00840	0,00840	19 383	0,02070	0,02070	21 775
2006 Q3	22 341	9 259	0,00900	0,85230	4 972	0,02260	1,11540	0,97500	21 782	0,00910	0,00910	19 209	0,02270	0,02270	21 580
2006 Q4	23 415	7 917	0,01240	0,70850	4 198	0,01740	1,01330	0,84730	19 840	0,01260	0,01260	17 496	0,01760	0,01760	19 656
2007 Q1	24 150	8 338	0,01240	0,74280	4 383	0,01480	1,13790	0,91930	22 201	0,01250	0,01250	19 579	0,01500	0,01500	21 996
2007 Q2	23 945	10 583	0,01930	0,97690	5 047	0,01770	1,45680	1,19290	28 564	0,01940	0,01940	25 190	0,01780	0,01780	28 299
2007 Q3	24 584	11 036	0,01930	1,03630	4 548	0,02230	1,45710	1,22880	30 209	0,01930	0,01930	26 640	0,02230	0,02230	29 928
2007 Q4	25 253	10 190	0,03490	0,97510	3 313	0,02350	1,25400	1,10580	27 924	0,03480	0,03480	24 625	0,02340	0,02340	27 664
2008 Q1	25 697	9 814	0,04290	1,00790	3 188	0,01380	1,52910	1,24150	31 903	0,04240	0,04240	28 133	0,01360	0,01360	31 606
2008 Q2	25 795	9 163	0,06240	1,05720	2 378	0,01920	1,36880	1,20290	31 028	0,06140	0,06140	27 364	0,01890	0,01890	30 741
2008 Q3	26 216	9 154	0,12660	1,27640	2 106	0,03090	1,66980	1,45990	38 272	0,12270	0,12270	33 751	0,03000	0,03000	37 917
2008 Q4	26 960	4 948	0,14700	1,24850	694	0,01720	1,49600	1,36670	36 847	0,14110	0,14110	32 493	0,01650	0,01650	36 504
Total	887 313	524 186			331 907				924 116	0,8339	0,8340	814 945	0,9709	0,9705	915 539

Table 6. Theoretical estimation of prior ultimate claims for the most recent accident quarters.

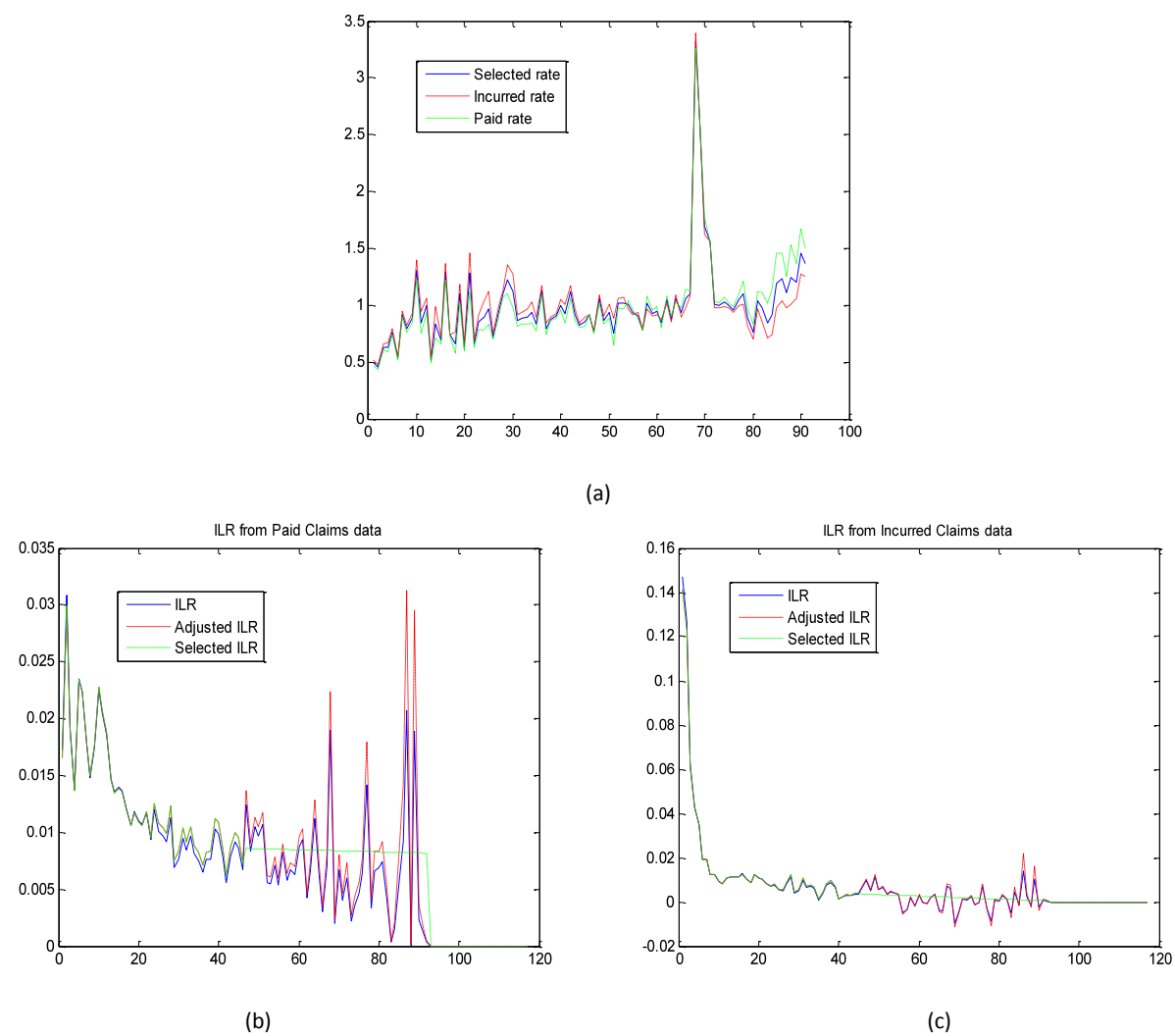


Figure 18. (a) – the on-level premium factors; (b) – the ILR from the paid claims data; (c) – the ILR from the incurred claims data. In all three graphs x-axis measures time in accident quarters (since 1986Q1).

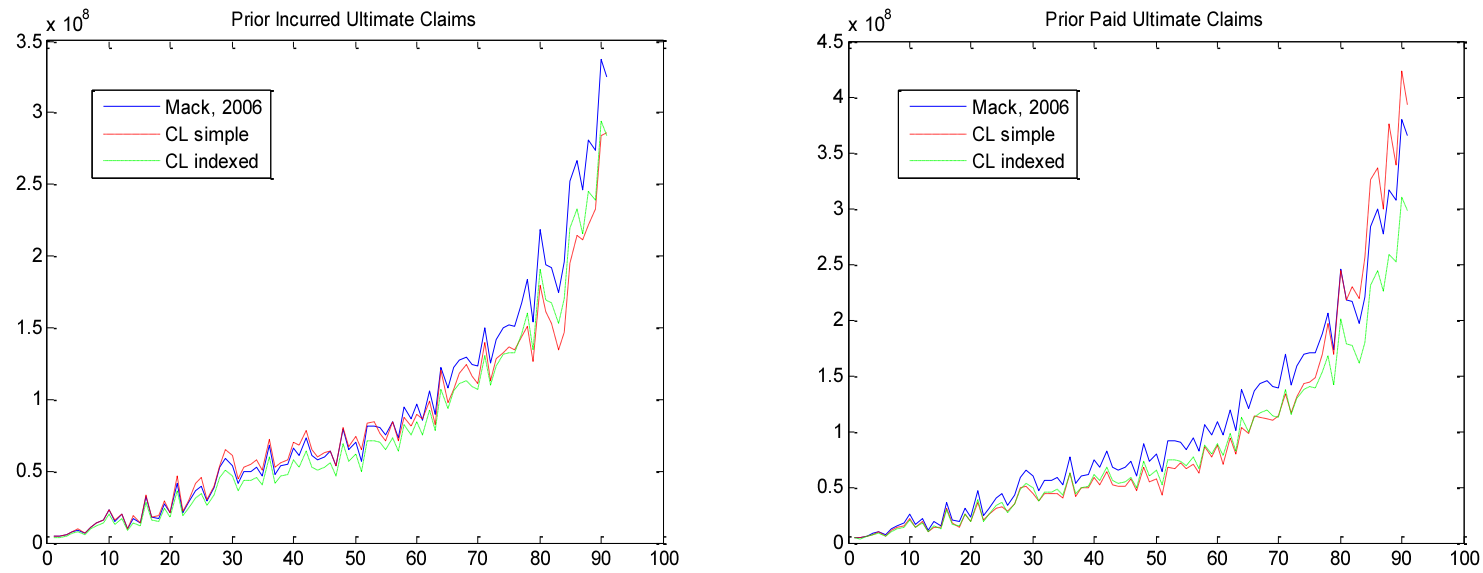


Figure 19. *The comparison of the three estimates of the prior ultimate claims for the incurred claims data (left graph) and the paid claims data (right graph). In both graphs x-axis measures time in accident quarters (since 1986Q1).*

The two practical estimations are based on the CL resulting ultimate claims. The first one just takes the resulting ultimate claims as prior estimate of ultimate loss. The second way assumes that the most representative quarters of the data are 2001 Q1 to 2002 Q4 due to most stable resulting ultimate claims. We take average ultimate loss of these quarters (see Table 7) and index it according to the on-level premium factors calculated above. The resulting prior ultimate loss ratios for the most recent accident quarters are given in the Table 8.

The three estimation results together are presented in the Figure 19 as time series of accident quarters. All three estimations provide with similar prior estimates; still, some important points must be noticed. For the paid claims theoretical estimation is the highest one until the accident period 2006 Q1, when the simple CL estimates becomes the highest. The indexed CL estimate for the paid claims is either approximately equal to the simple CL estimate, or the lowest one. For the incurred claims we see that until 2002 Q1 accident period all three estimations are almost equal. In the later accident periods the theoretical estimation gives the highest estimate, while the two CL estimates are very similar until year 2005, when the indexed CL prior ultimate claims become clearly higher than the simple CL.

Average Incurred LR 2001 Q1 - 2002 Q4	81,00%	Average Paid LR 2001 Q1 - 2002 Q4	76,97%
--	---------------	--	---------------

Table 7. The average ULR from the accident quarters 2001 Q1 – 2002 Q4.

Acc. Qrt	CL indexed		CL simple	
	Prior Ultimate LR (incurred claims)	Prior Ultimate LR (paid claims)	Prior Ultimate LR (incurred claims)	Prior Ultimate LR (paid claims)
2005 Q1	77,70%	73,83%	74,92%	82,44%
2005 Q2	85,07%	80,83%	79,60%	93,65%
2005 Q3	89,35%	84,91%	79,95%	104,21%
2005 Q4	71,80%	68,23%	64,05%	85,76%
2006 Q1	61,71%	58,64%	55,25%	75,18%
2006 Q2	84,43%	80,23%	76,60%	103,10%
2006 Q3	78,98%	75,05%	68,27%	102,67%
2006 Q4	68,63%	65,22%	57,46%	93,84%
2007 Q1	74,46%	70,76%	60,83%	105,96%
2007 Q2	96,63%	91,82%	81,29%	136,25%
2007 Q3	99,53%	94,58%	87,01%	136,99%
2007 Q4	89,57%	85,11%	83,55%	118,69%
2008 Q1	100,56%	95,56%	86,31%	146,08%
2008 Q2	97,44%	92,59%	90,26%	131,28%
2008 Q3	118,25%	112,37%	108,37%	161,26%
2008 Q4	110,70%	105,20%	105,98%	145,73%

Table 8. The two practical estimations' resulting prior ultimate loss ratios.

4.2.2.2 *Estimation of other BF parameters*

Having the prior ultimate claims estimated, the estimates for the other two unknown parameters, y_k and the s_k^2 , will also be calculated. The resulting estimations are presented graphically in the Figures 20 and 21, correspondingly for incurred and paid claims data. According to the procedure, explained earlier in the thesis, the following steps have been made:

- The raw estimates of y_k :s were calculated (upper left graph)
- The linear regression on $\ln(y_k)$ was fitted (upper right graph)
- The proportionality constants s_k^2 were calculated for $k = 1, \dots, n - 1$, and s_n^2, s_{n+1}^2 were extrapolated (lower left graph)
- The minimization of the function Q (10) was implemented and the final values of y_k obtained (upper left graph)
- The development pattern $\hat{z}_k^* = \hat{y}_1^* + \dots + \hat{y}_k^*$ was calculated and the remaining percentage to develop per accident quarter obtained (lower right graph)

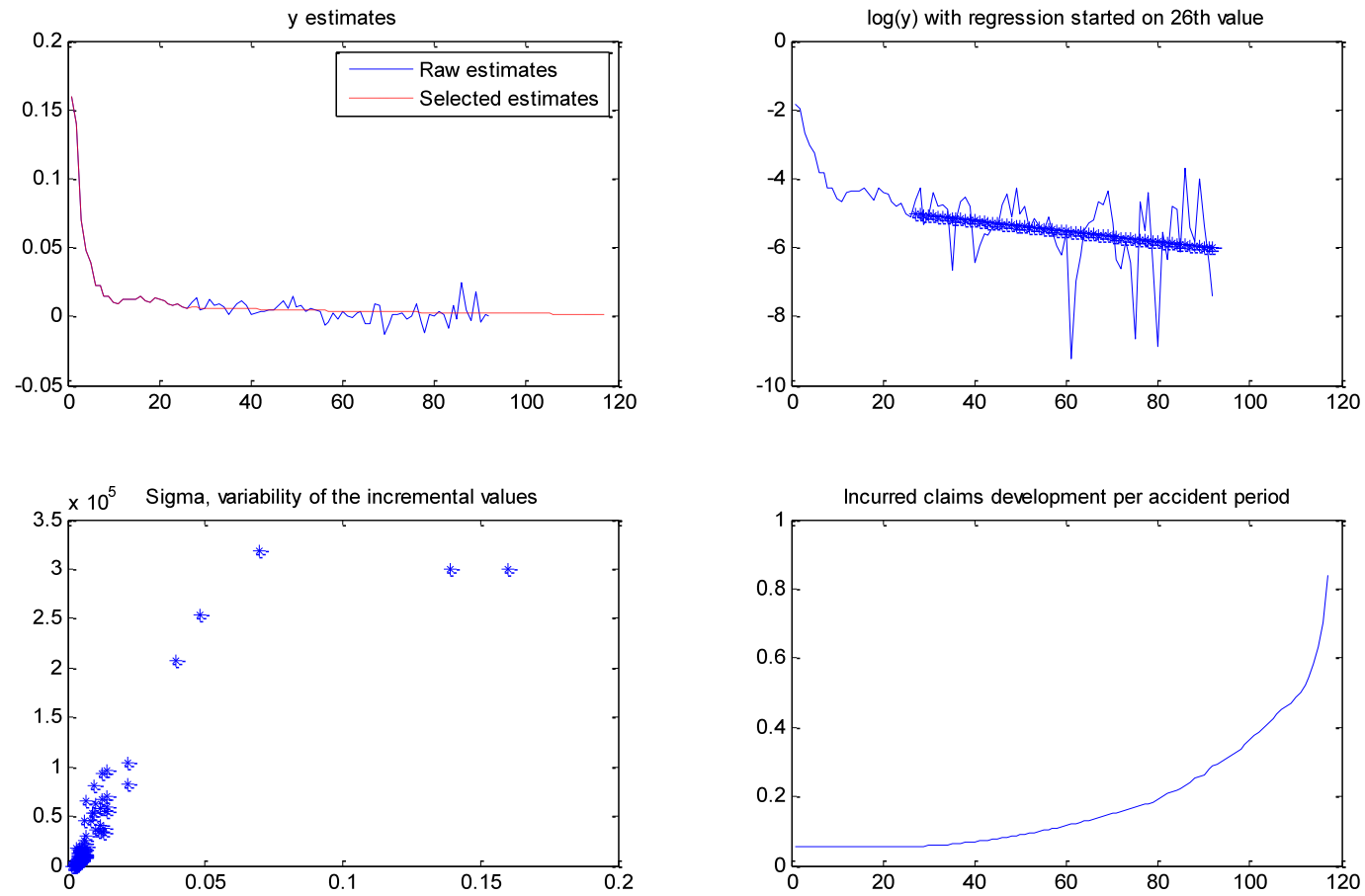


Figure 20. *The estimation of other BF parameters from the incurred claims data. In all the graphs x-axis measures time in accident quarter (since 1986Q1).*

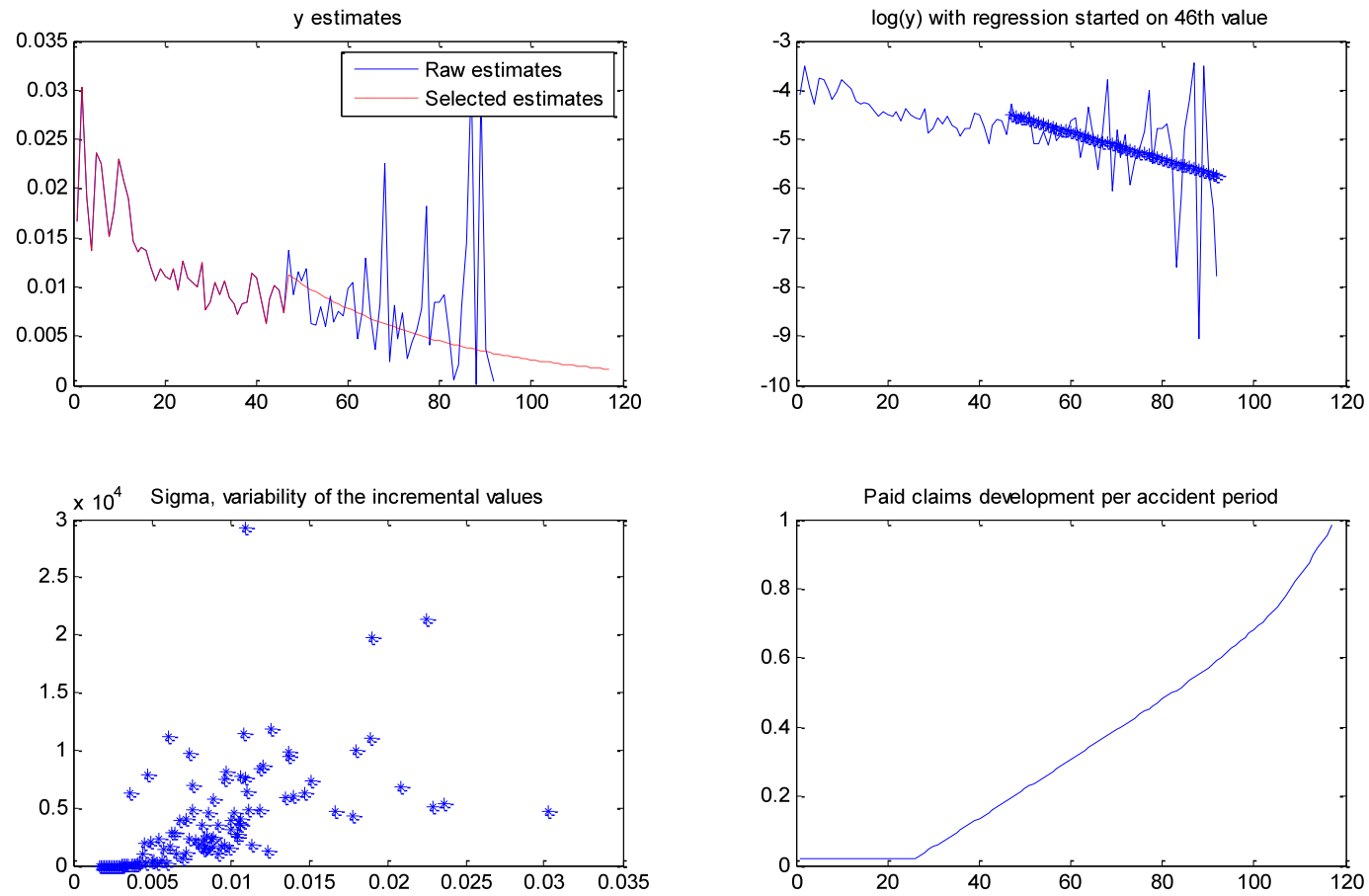


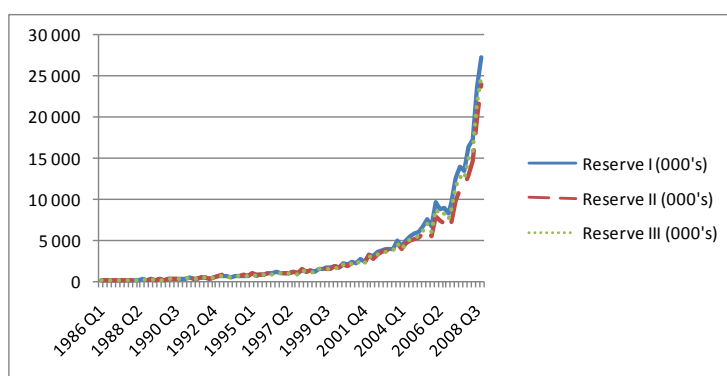
Figure 21. *The estimation of other BF parameters from the paid claims data. In all the graphs x-axis measures time in accident quarters (since 1986Q1).*

4.2.2.3 Theoretical Prediction Error

Having all the parameters estimated we can now present the BF reserve. Three different reserve estimates will be presented corresponding to the three ways of estimating the prior ultimate claims. In the Tables 9-10 below the most recent accident quarters' results are given, and graphically time series of the reserve estimates are presented since 1986 Q1 (see Figures 22-23). In the following *Reserve I* will be the estimate of the BF IBNR reserve when the prior ultimate claims are according to Mack (2006); *Reserve II* – when the prior ultimate claims are the simple CL ultimate claims; *Reserve III* – the indexed CL ultimate claims.

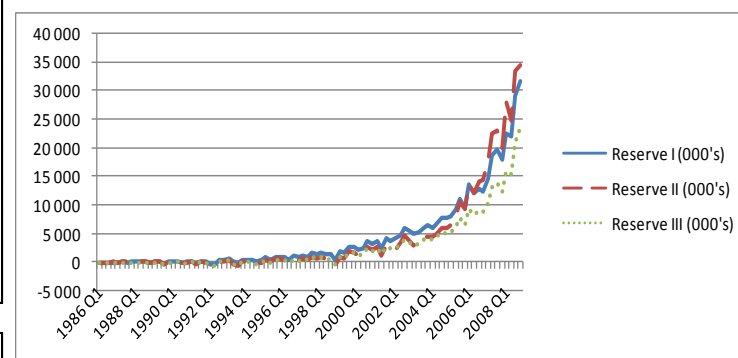
Acc. Qrt	Reserve I (000's)	Reserve II (000's)	Reserve III (000's)
2005 Q1	5 851	5 182	5 374
2005 Q2	6 689	5 750	6 144
2005 Q3	7 577	6 228	6 960
2005 Q4	6 550	5 367	6 017
2006 Q1	9 584	7 882	8 804
2006 Q2	8 752	7 293	8 038
2006 Q3	8 856	7 031	8 134
2006 Q4	8 247	6 341	7 574
2007 Q1	9 507	7 133	8 732
2007 Q2	12 590	9 729	11 564
2007 Q3	13 901	11 161	12 768
2007 Q4	13 389	11 471	12 298
2008 Q1	16 406	12 934	15 068
2008 Q2	17 272	14 696	15 864
2008 Q3	23 652	19 910	21 725
2008 Q4	27 293	24 001	25 071
Total	292 424	255 001	268 597

Table 9, Figure 22. *The BF reserve estimate based on the Incurred Claims data*



Acc. Qrt	Reserve I (000's)	Reserve II (000's)	Reserve III (000's)
2005 Q1	7 950	6 362	5 276
2005 Q2	9 342	7 996	6 336
2005 Q3	11 123	10 423	7 760
2005 Q4	9 492	9 191	6 621
2006 Q1	13 619	13 547	9 460
2006 Q2	12 114	12 087	8 328
2006 Q3	12 981	14 068	9 130
2006 Q4	12 461	14 369	8 852
2007 Q1	14 543	17 565	10 415
2007 Q2	18 691	22 394	13 286
2007 Q3	19 673	22 951	13 836
2007 Q4	17 928	19 998	12 394
2008 Q1	22 459	27 918	15 970
2008 Q2	21 928	24 844	15 521
2008 Q3	29 090	33 245	21 029
2008 Q4	31 641	34 379	23 634
Total	391 914	369 190	261 589

Table 10, Figure 23. *The BF reserve estimate based on the Paid Claims data*



Finally, the prediction error for the reserve estimates of the BF reserve was calculated. In the formulae presented earlier in Theorem 2 and Corollary 2 one can estimate everything except $s.e.(\hat{U}_i)$ because the prior ultimate claims are correlated. The solution was borrowed from Mack (2008) where the author computes coefficient of variation

$$c.v.(\hat{U}_i) = \frac{s.e.(\hat{U}_i)}{\hat{U}_i} \approx \frac{\sqrt{(s.e.(y_1))^2 + \dots + (s.e.(y_{n+1}))^2}}{y_1 + \dots + y_{n+1}}$$

and then adds a subjectively estimated variability of on-level premium factor and premiums itself.

For our data the following coefficients of variation were obtained:

$$c.v.(\hat{U}_i)_{incurred}^I = 3,77\%, c.v.(\hat{U}_i)_{paid}^I = 2,30\%$$

$$c.v.(\hat{U}_i)_{incurred}^{II} = 4,86\%, c.v.(\hat{U}_i)_{paid}^{II} = 3,53\%$$

$$c.v.(\hat{U}_i)_{incurred}^{III} = 4,35\%, c.v.(\hat{U}_i)_{paid}^{III} = 2,82\%$$

To be able to compare the prediction error of the reserve estimate from paid claims and incurred claims data, as well as among the three types of estimation, it was decided to calculate the prediction error with $s.e.(\hat{U}_i) = 6,6\%$ and $s.e.(\hat{U}_i) = 10\%$. This means, that, for example, for the reserve estimate of type I from the incurred claims data, 2,83% of variability is added subjectively with $s.e.(\hat{U}_i) = 6,6\%$, and 6,23% of variability is added subjectively with $s.e.(\hat{U}_i) = 10\%$. It appears that the difference in $s.e.(\hat{U}_i)$ creates minor variability in the resulting prediction errors, the higher the assumed $s.e.(\hat{U}_i)$, the higher the prediction error, of course.

Correspondingly how the reserve estimates were denoted above, we denote *Prediction Error I*, *Prediction Error II* and *Prediction Error III*, representing the type of prior ultimate claims estimation. The Figures 24 and 25 below confirms that the difference between $s.e.(\hat{U}_i) = 6,6\%$ and $s.e.(\hat{U}_i) = 10\%$ is only minor. The difference among the three types of prediction errors is more significant. For the paid claims *Prediction Error I* and *III* are almost equal, *Prediction Error II* being the highest. As for the incurred claims *Prediction Error II* is also the highest one, but *Prediction Error III* is higher than the *Prediction Error I*. The full time series of all prediction errors since accident quarter 1986 Q1 can be found in the Figure 26 for the incurred claims data, and the Figure 27 for the paid claims data. The most recent accident quarters' results are presented in the Tables 11 for the incurred claims data and 12 for the paid claims data.

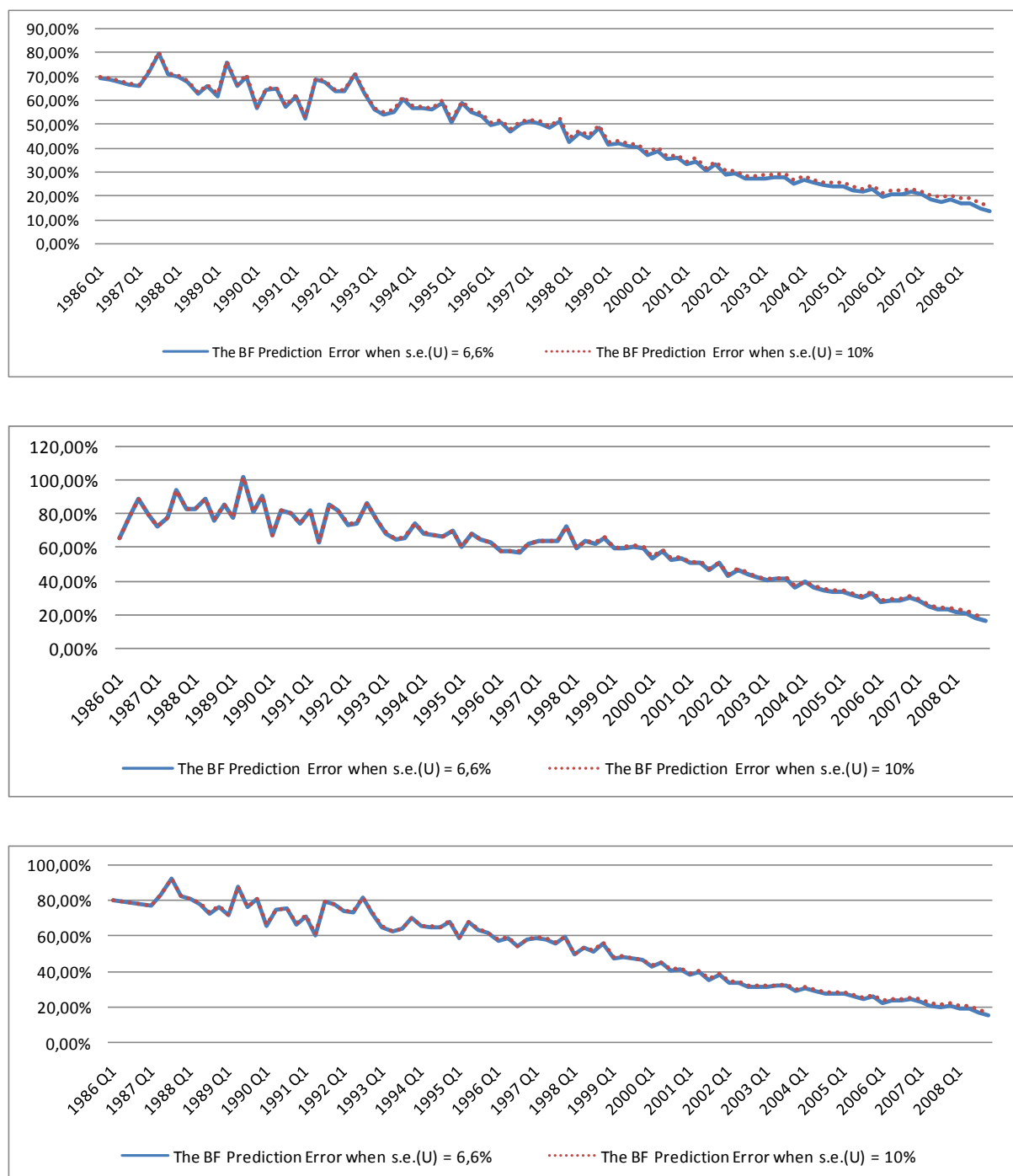


Figure 24. The prediction error of the BF reserve estimate from incurred claims data (I – the upperst graph, II – the middle graph, III – the lowest graph).

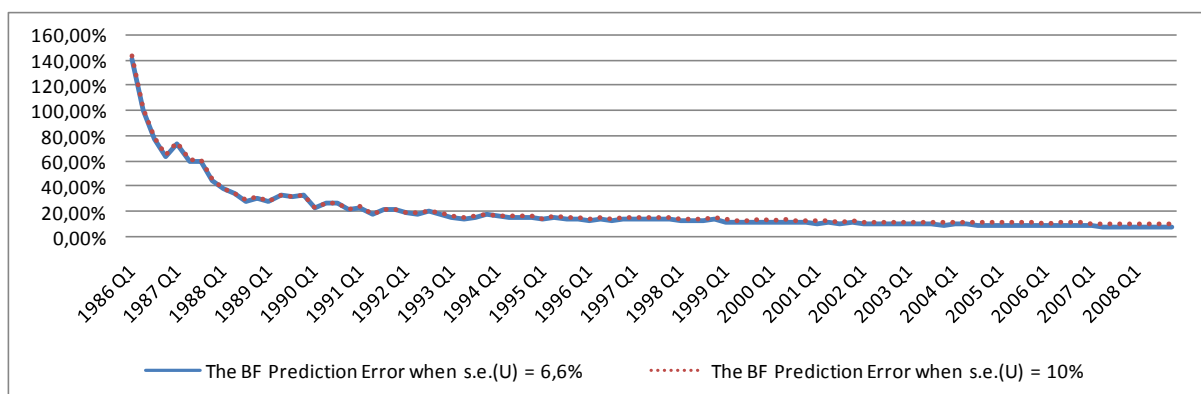
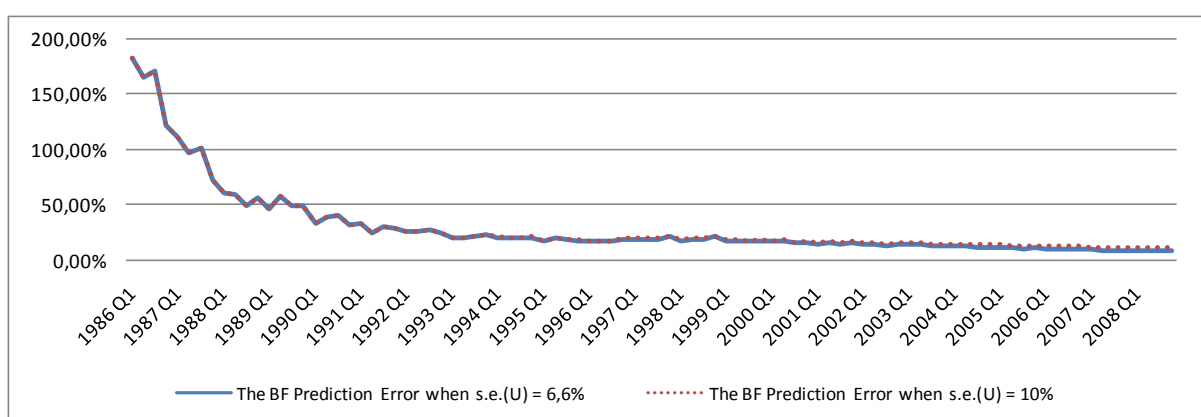
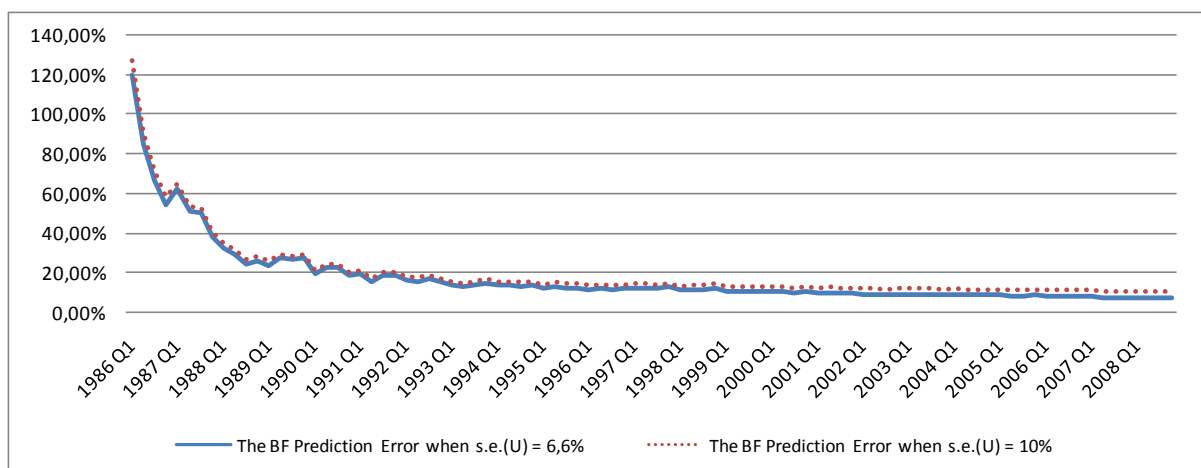


Figure 25. The prediction error of the BF reserve estimate from paid claims data (I – the upper graph, II – the middle graph, III – the lowest graph).

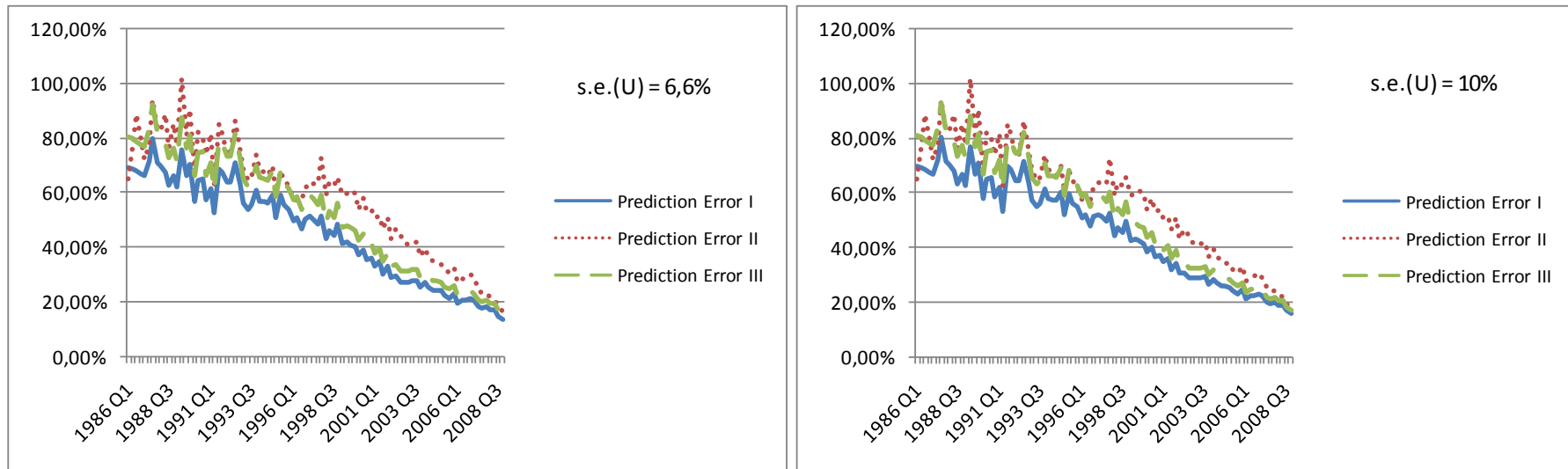


Figure 26. Prediction Error I, II and III for the incurred claims data.

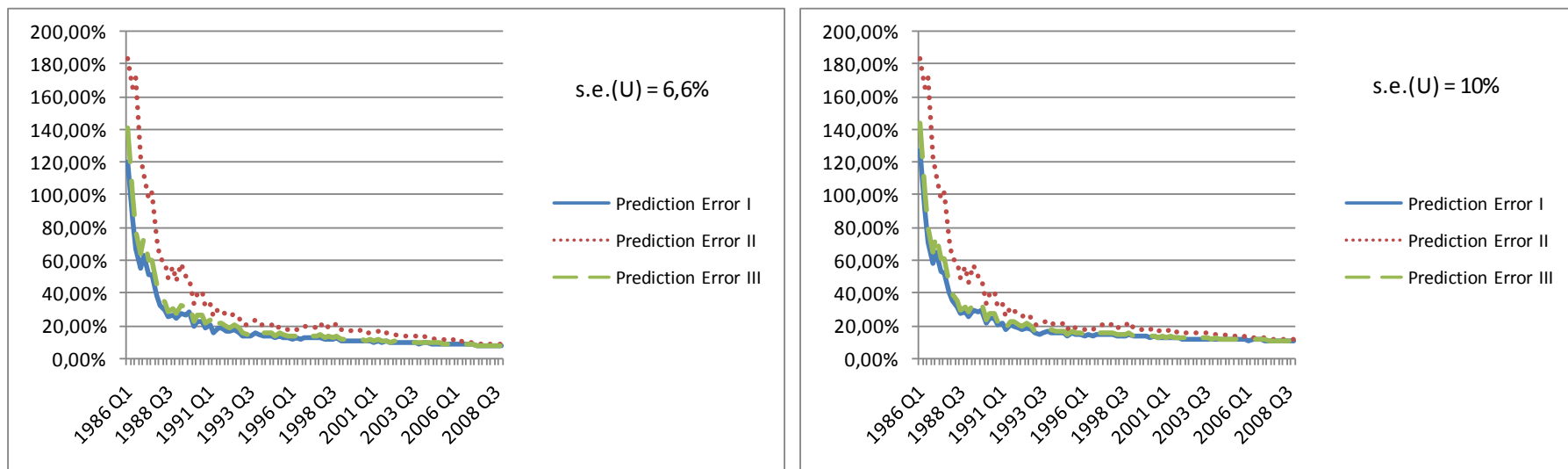


Figure 27. The prediction Error I, II and III for the paid claims data.

Acc. Qrt	s.e.(U) = 6,6%			s.e.(U) = 10%		
	Prediction Error I (%)	Prediction Error II (%)	Prediction Error III (%)	Prediction Error I (%)	Prediction Error II (%)	Prediction Error III (%)
2005 Q1	24,12%	33,52%	27,54%	25,26%	34,19%	28,67%
2005 Q2	22,62%	31,58%	25,81%	23,84%	32,29%	27,00%
2005 Q3	21,60%	30,35%	24,63%	22,87%	31,10%	25,87%
2005 Q4	22,97%	32,31%	26,20%	24,17%	33,02%	27,37%
2006 Q1	19,77%	27,10%	22,50%	21,15%	27,94%	23,84%
2006 Q2	20,77%	28,11%	23,65%	22,09%	28,93%	24,92%
2006 Q3	20,83%	28,64%	23,71%	22,14%	29,45%	24,98%
2006 Q4	21,59%	30,03%	24,59%	22,86%	30,81%	25,80%
2007 Q1	20,51%	28,49%	23,34%	21,85%	29,32%	24,61%
2007 Q2	18,30%	24,65%	20,78%	19,79%	25,60%	22,20%
2007 Q3	17,62%	23,05%	19,99%	19,16%	24,08%	21,45%
2007 Q4	18,34%	23,01%	20,83%	19,82%	24,06%	22,22%
2008 Q1	17,10%	21,80%	19,40%	18,68%	22,92%	20,87%
2008 Q2	17,02%	20,64%	19,38%	18,61%	21,85%	20,85%
2008 Q3	14,94%	18,14%	16,99%	16,72%	19,54%	18,63%
2008 Q4	13,63%	16,23%	15,47%	15,56%	17,84%	17,22%
Total	6,25%	7,41%	7,11%	7,03%	7,91%	7,80%

Table 11. The prediction Error for the Incurred Claims data.

Acc. Qrt	s.e.(U) = 6,6%			s.e.(U) = 10%		
	Prediction Error I (%)	Prediction Error II (%)	Prediction Error III (%)	Prediction Error I (%)	Prediction Error II (%)	Prediction Error III (%)
2005 Q1	8,76%	11,30%	9,56%	11,54%	13,30%	12,04%
2005 Q2	8,55%	10,72%	9,28%	11,38%	12,83%	11,83%
2005 Q3	8,37%	10,14%	9,03%	11,25%	12,38%	11,64%
2005 Q4	8,64%	10,57%	9,41%	11,45%	12,75%	11,95%
2006 Q1	8,09%	9,43%	8,65%	11,04%	11,85%	11,37%
2006 Q2	8,21%	9,66%	8,84%	11,13%	12,06%	11,52%
2006 Q3	8,17%	9,43%	8,80%	11,10%	11,91%	11,50%
2006 Q4	8,24%	9,44%	8,91%	11,15%	11,94%	11,59%
2007 Q1	8,05%	9,03%	8,65%	11,01%	11,64%	11,40%
2007 Q2	7,75%	8,52%	8,23%	10,79%	11,26%	11,09%
2007 Q3	7,71%	8,48%	8,18%	10,76%	11,25%	11,06%
2007 Q4	7,75%	8,63%	8,25%	10,79%	11,38%	11,12%
2008 Q1	7,59%	8,21%	8,03%	10,68%	11,09%	10,97%
2008 Q2	7,63%	8,37%	8,10%	10,71%	11,21%	11,03%
2008 Q3	7,43%	8,01%	7,81%	10,56%	10,96%	10,82%
2008 Q4	7,42%	8,04%	7,81%	10,56%	11,00%	10,83%
Total	2,67%	2,95%	2,77%	3,87%	4,15%	3,94%

Table 12. The prediction Error for the Paid Claims data.

4.2.2.4 Bootstrapped Estimation Error

The estimation error of the BF reserve was bootstrapped according to the procedure described in the section 3.2.4. Taking one prior ultimate claims estimation method at a time, two bootstrapping approaches explained earlier were implemented. The two bootstrap approaches were then compared to the theoretical estimation error. In the Figure 28 (a)-(c) the estimation error for the incurred claims data is presented; and in the Figure 29 (a)-(c) the same for the paid claims data is shown.

It is easy to notice that the theoretical estimation error lies in-between the two bootstrap approaches, the second approach providing with the highest error, and the first approach, when the prior ultimate claims are kept as a constant, gives the lowest error.

For the incurred claims data the estimation method of the prior ultimate claims does not play a significant role for the size of the estimation error. What is more interesting, the second bootstrap approach has exactly the opposite trend of the estimation error in comparison to the theoretical one.

For the paid claims data, the choice of the estimation method affects the second bootstrapping approach, giving lowest estimation error with the *method I*, and highest with the *method II*. Here, the theoretical estimation error is almost a constant throughout all the accident quarters. The two bootstrap approaches have peak of the estimation error during the older accident quarters, then the error decreases during the middle accident periods. For the most recent accident quarters the first bootstrapping approach error continues to decrease, while the second bootstrapping approach error increases.

The second bootstrap approach provides us with extremely high reserve estimate. Graphically this can be seen in the figures in Appendix IV. The phenomenon appears due to too volatile prior ultimate claims estimates created in each bootstrap loop. Therefore, we decide not to believe in the result of the second bootstrap approach.

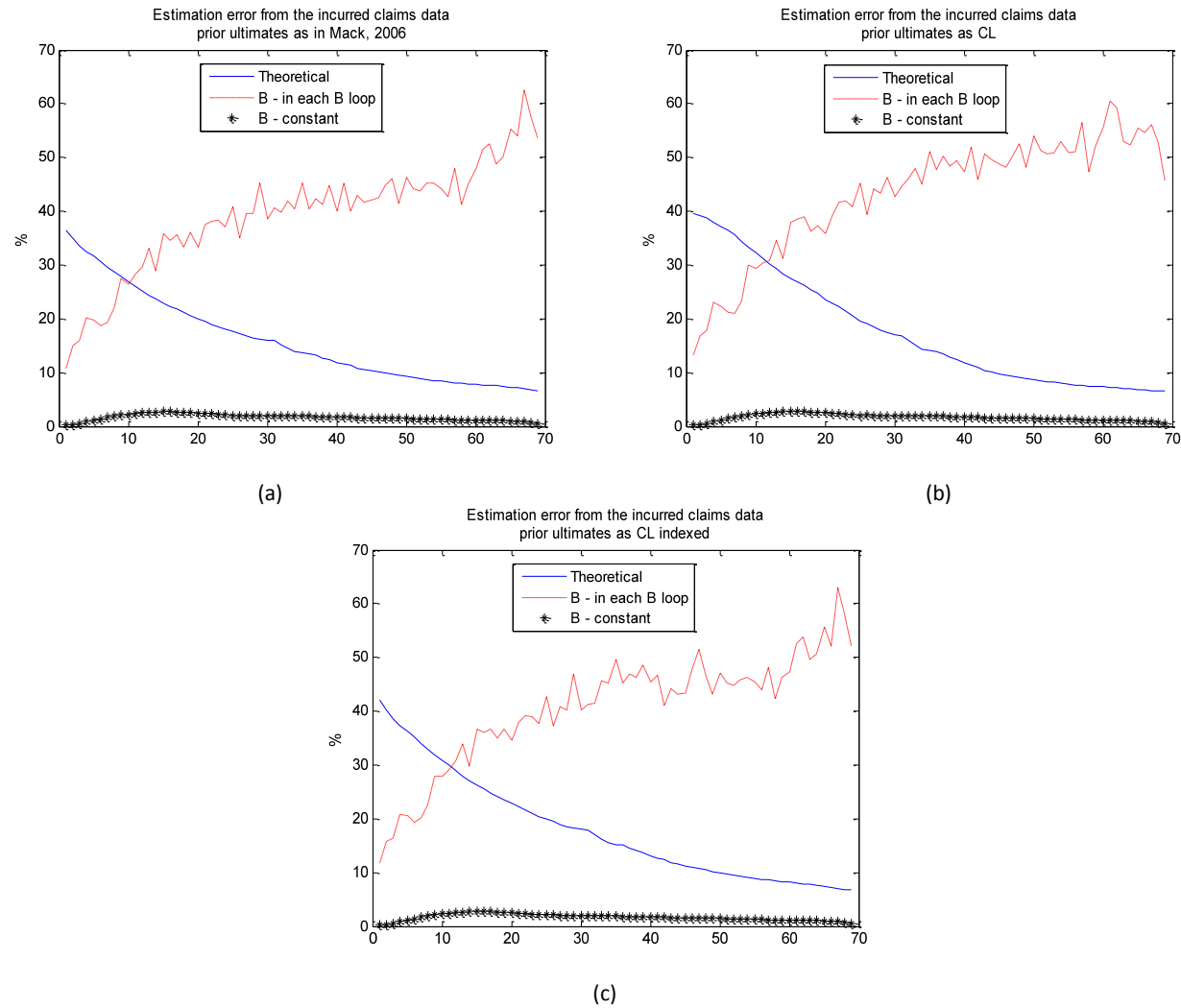


Figure 28. The estimation error of the BF reserve, when the prior ultimate claims as in Mack (2006) (a), as CL ultimates (b), as CL indexed ultimates (c); incurred claims data. In all the graphs x-axis measures time in accident quarters (since 1991Q1).

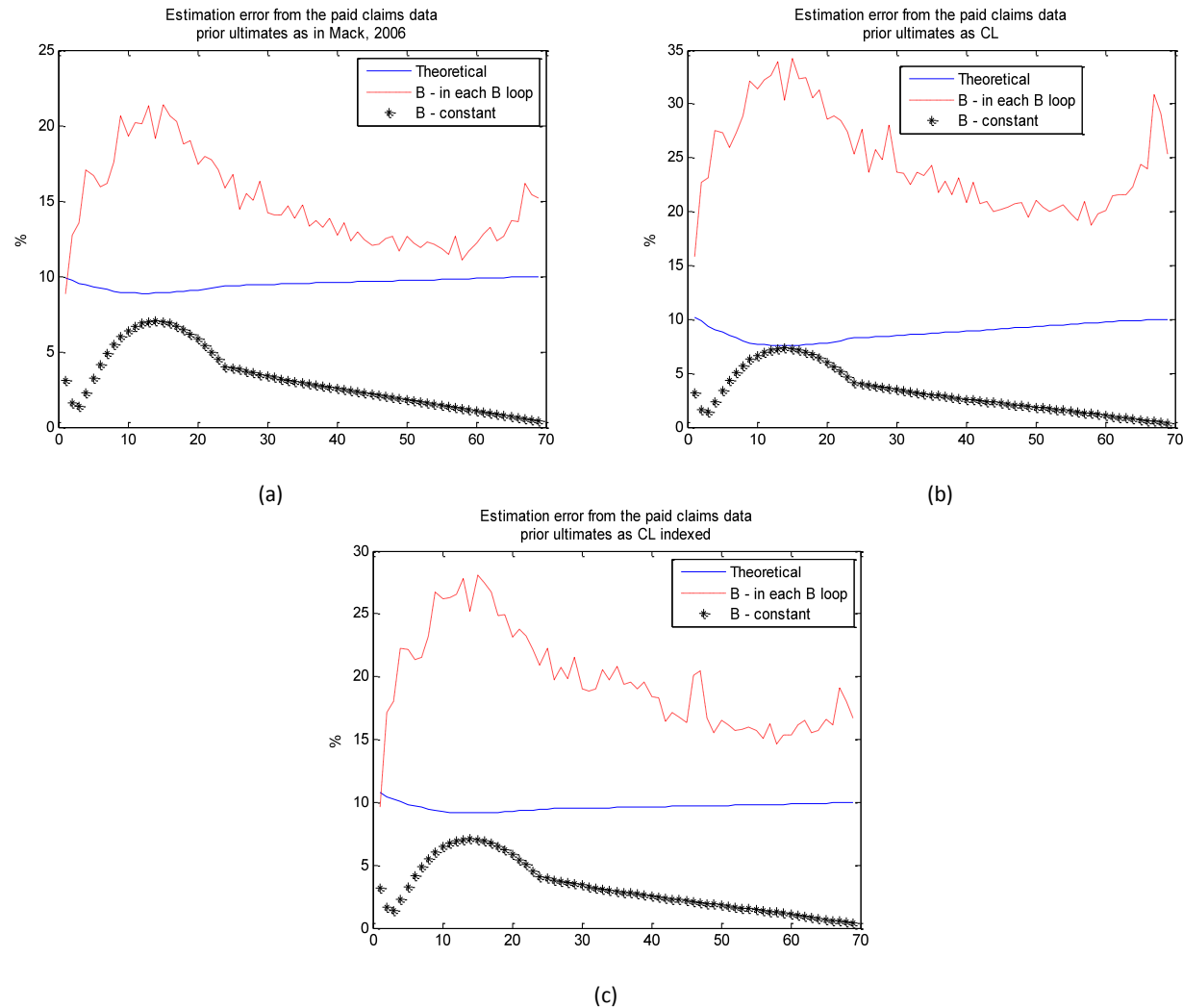


Figure 29. The estimation error of the BF reserve, when the prior ultimate claims as in Mack (2006) (a), as CL ultimates (b), as CL indexed ultimates (c); paid claims data. In all the graphs x-axis measures time in accident quarters (since 1991Q1).

4.2.3 Chain-Ladder vs Bornhuetter-Ferguson

4.2.3.1 Theoretical Prediction Error

The two reserving methods being available and very widely used, actuaries are interested which of the two gives smaller prediction error. As a prejudice, at least for the most recent quarters one would expect to get a smaller prediction error with the BF method. But what happens during all other accident periods and with the total reserve prediction error? In the Figures 30-32 the graphical comparison of the CL Prediction Error with all three types of the BF prediction error is presented. In the Table 13 you find the total reserve estimates' prediction errors for all the methods used above.

We clearly see that the CL method has a higher prediction error for the total reserve estimate, as well as for almost all the accident quarters' estimates, does not matter which BF estimation we choose. Even when the BF method is using simple CL ultimate claims as its prior estimates, the prediction error is significantly smaller than the CL prediction error. The best result, giving the smallest prediction error, was obtained with the theoretical estimation of prior ultimate claims (BF I), which could have been expected, since the prior estimates do not depend on the latest diagonal of the data. Moreover, the CL prediction error increases for the most recent accident quarters due to the lack of data, meanwhile the BF prediction error keeps stable in these particular accident quarters.

A natural question arises: why would anybody use the CL method, if the BF method provides with more exact results in any case? One possible reason could be the prior belief not always depicting the best level of ultimate claims. If the belief fits well with what has been happening in the business recently, it does not necessarily mean that the same can be applied for the historical data. Here the CL method, which is based on the historical data, comes to help us. For older accident periods, where we have rather much data observed, it might be much better to use the CL method with a higher prediction error, than the BF method, based on a non-representative belief.

	CL	BF I	BF II	BF III
Total Reserve Prediction Error (incurred claims)	11,53%	7,03%	7,91%	7,80%
Total Reserve Prediction Error (paid claims)	6,12%	3,87%	4,15%	3,94%

Table 13. *The prediction errors of the total reserve for all the methods considered.*

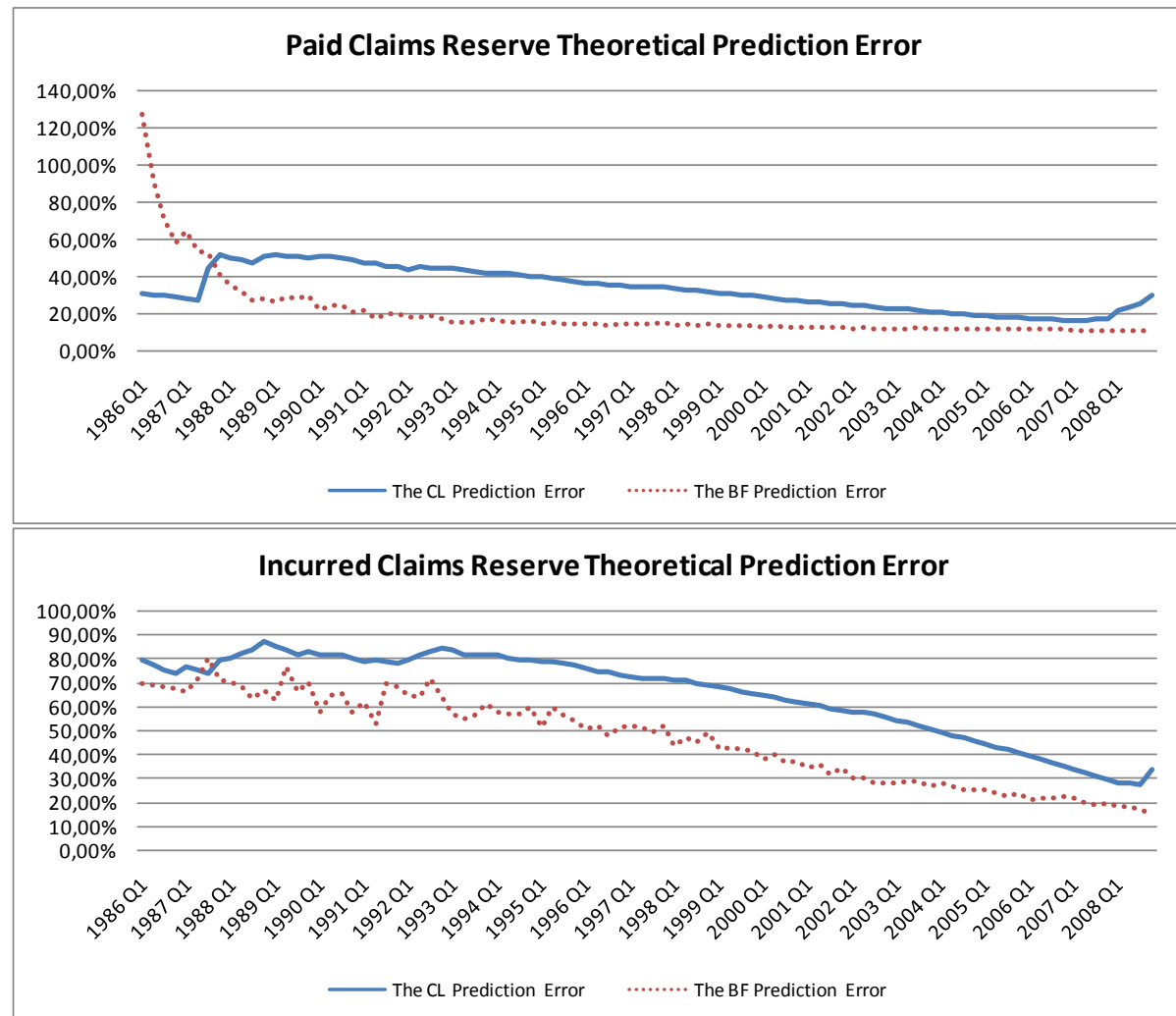


Figure 30. The comparison of the Prediction Error I with the CL Prediction Error for both, paid and incurred claims data.

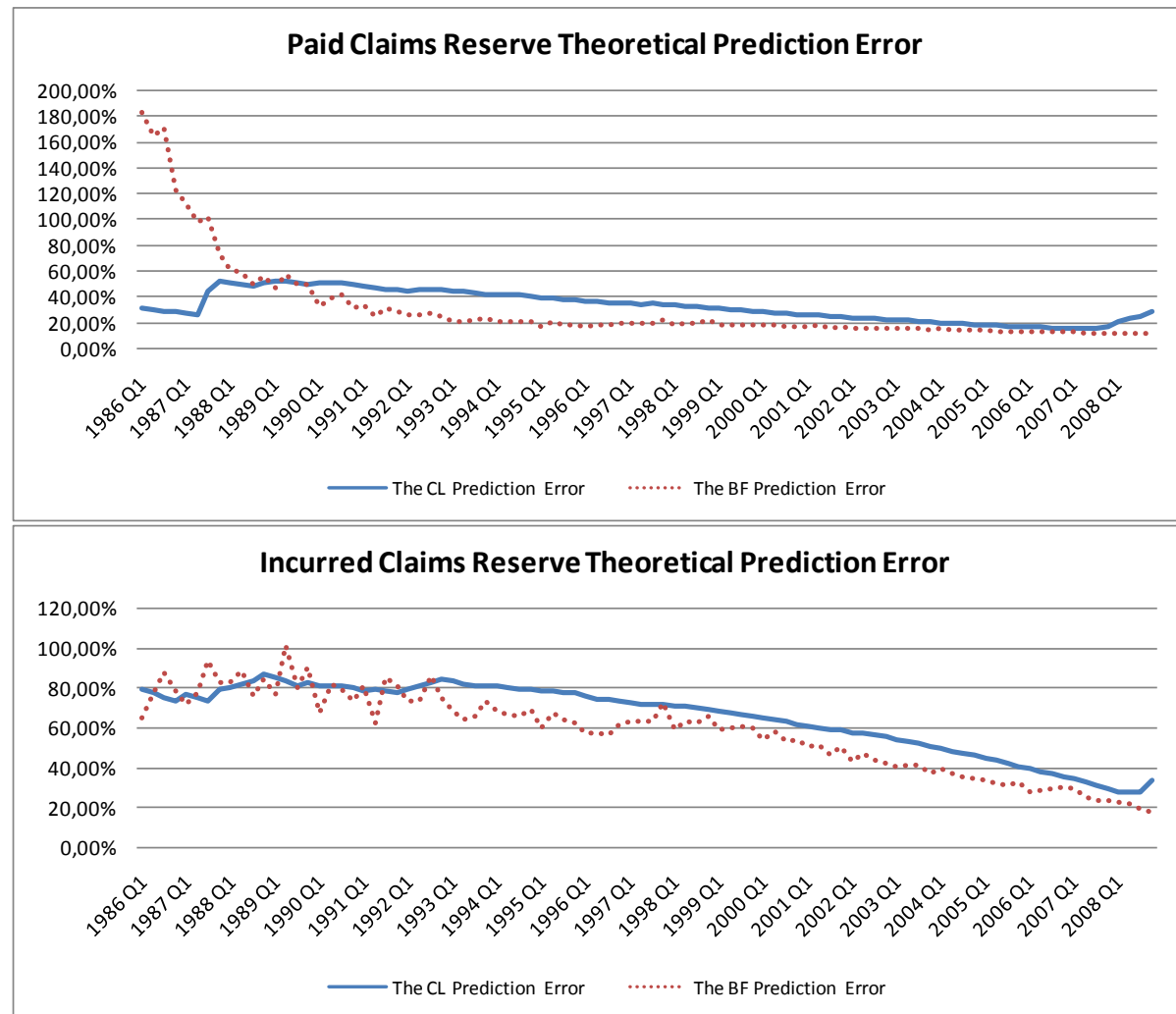


Figure 31. The comparison of the Prediction Error II with the CL Prediction Error for both, paid and incurred claims data.

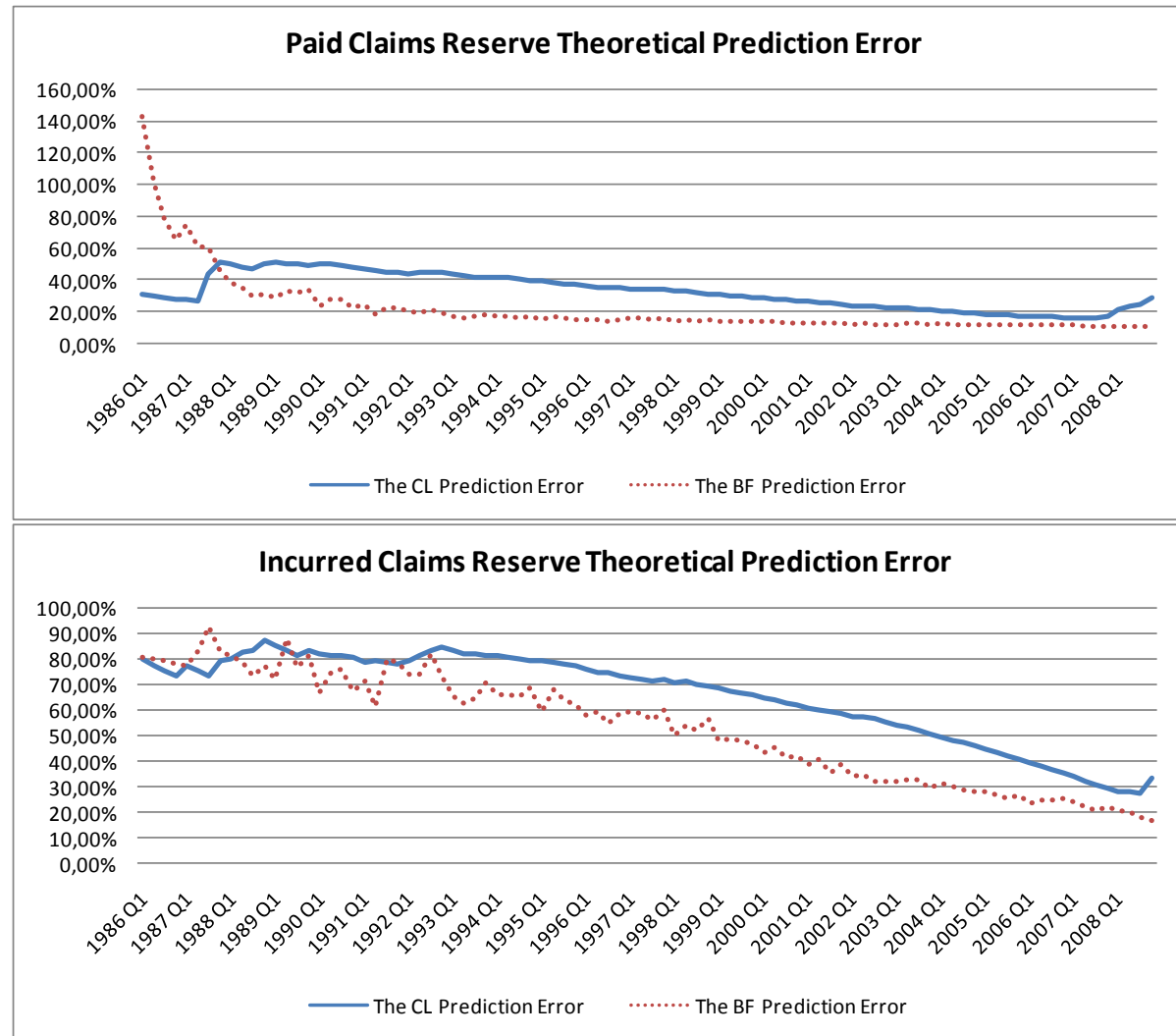


Figure 32. The comparison of the Prediction Error III with the CL Prediction Error for both, paid and incurred claims data.

4.2.3.2 Bootstrapped Estimation Error

When it comes to bootstrapping, only the estimation error of the two methods can be compared. In the following the histograms of the bootstrap total reserve estimates will be discussed. In the Figures 16-17 (a) the CL method estimation error histograms were presented for the incurred claims and the paid claims data. Below, in the Figures 33 and 34 the BF method estimation error histograms are shown. The *Estimation Error I* represents the prior ultimate estimation I, the *Estimation Error II* – the prior ultimate estimation II, and the *Estimation Error III* – the estimation III.

First, the size of the mean of the bootstrap reserve estimate must be noticed. For the incurred claims the Bootstrap Reserve Estimate I and the Bootstrap Reserve Estimate II are approximately the same size as the CL Bootstrap Reserve Estimate (see Figure 16 (a)). For the paid claims data, the Bootstrap Reserve Estimate I is almost exactly equal to the CL one (see Figure 17 (a)), and the Bootstrap Reserve Estimate II is little higher than the CL Bootstrap Reserve Estimate. The Bootstrap Reserve Estimate III is much lower than the CL Bootstrap Reserve Estimate for both, incurred and paid claims.

As for the estimation error itself, or the spread of the histograms, the BF estimation error generally is smaller than the CL estimation error. For the incurred claims data the difference between the smallest and the largest bootstrap estimates is approximately $diff \approx 0.35 * 10^9$, and does not depend on the choice of the prior ultimate estimation method. For the paid claims data the choice of the prior ultimate estimation method does not play a significant role either. The *Estimation Error I* is the highest one, but the difference $diff \approx 0.8 * 10^9$ is approximately the same for all three types of estimation error. For comparison the difference between the smallest and the largest bootstrap estimates of the CL method was approximately $2.5 * 10^9$ for both, incurred and paid claims data.

Notice, that the estimation error of the BF Bootstrap Reserve Estimate is clearly higher for the paid claims data, which was not seen for the CL method. What is more, the theoretical prediction error for both methods was smaller for the paid claims data. The best guess here could be that the process error of the BF method is much higher for the incurred claims data. Still, no real conclusion can be made without further investigation.

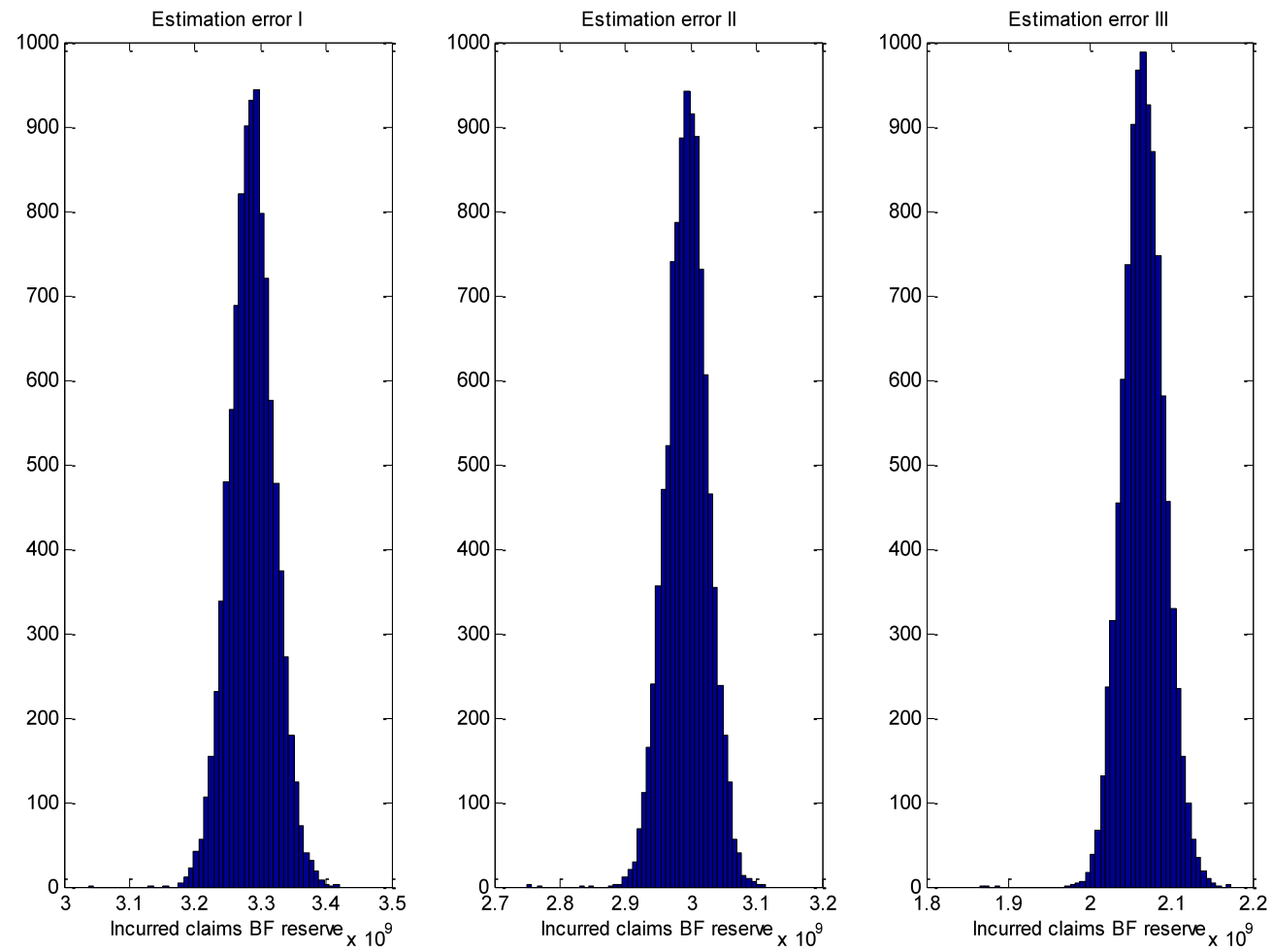


Figure 33. *The BF bootstrapped estimation error corresponding to the three prior ultimate claims estimation types; incurred claims data.*

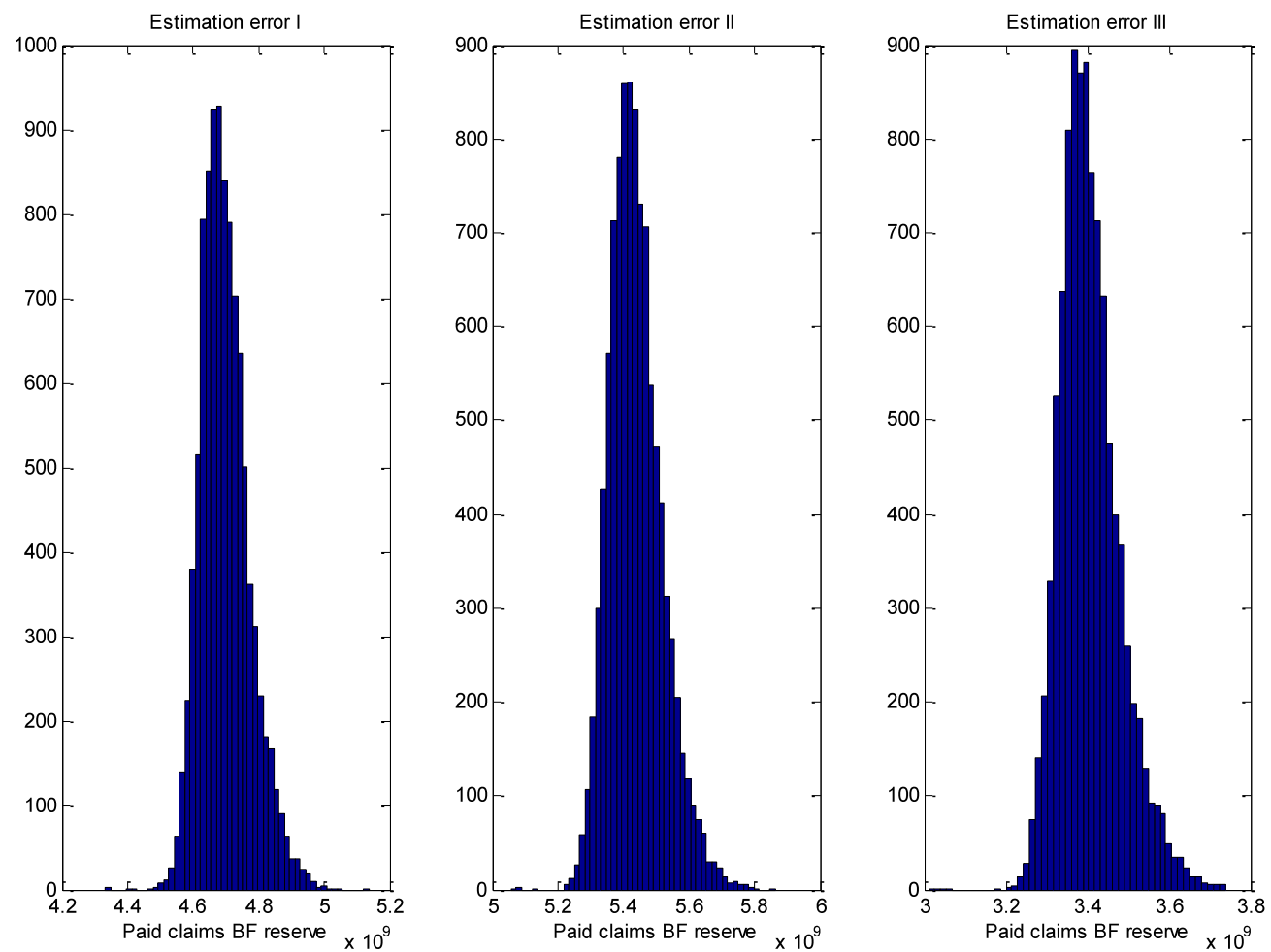


Figure 34. The BF bootstrapped estimation error corresponding to the three prior ultimate claims estimation types; paid claims data.

5. Conclusions

After the numerical study has been performed the following statements are the most important:

1. The prediction error for the paid claims data is smaller than for the incurred claims data.
2. The CL method prediction error is significantly higher than the BF method prediction error.
3. The CL bootstrapped estimation error is much higher than the BF bootstrapped estimation error, does not matter which way of prior ultimate estimation is chosen.
4. The CL bootstrapped prediction error is mostly lower than the theoretical CL prediction error for separate accident periods. The bootstrap technique allows for higher covariance between incurred data than theoretical calculation.
5. The CL method estimation error is higher than the process error.
6. The CL prediction error decreases in time, and suddenly increases for the most recent accident periods due to the lack of data.
7. The BF prediction error decreases in time.
8. The Mack (2006) prior ultimate claims calculation gave the lowest theoretical prediction error. On the other hand, the prior ultimate estimation methods did not have a significant impact on the bootstrapped estimation error.
9. The BF bootstrapped estimation error is much higher than the theoretical one, when the prior ultimate claims are estimated in each bootstrap loop from pseudo-data. In this way the reserve estimate is too high, therefore, the method is not seen as a good one.
10. The BF bootstrapped estimation error is much lower than the theoretical one, when the prior ultimate claims are kept as a constant throughout the bootstrapping loops.

The above is true for the Personal Accident Children insurance data at Trygg-Hansa.

The lower prediction error for the paid claims data or for the BF method does not simultaneously mean that the method is better to use. The actuary has to judge subjectively the reliability of the data and prior beliefs of the ultimate claims. Taking all the surrounding knowledge and the prediction error into account the decision must be made.

The most significant impact into the actuarial mathematics is the attempt to bootstrap the BF method estimation error. The process error, and therefore the prediction error, of the BF method were not analyzed in the bootstrap procedure. This could be the case for further investigation.

6. References

- [1] Alai, D. H., Merz, M., Wütrich, M.V., *Mean Square Error of Prediction in the Bornhuetter-Ferguson Claims Reserving Method* (2008), Department of Mathematics, ETH Zurich, CH-8092 Zurich, Switzerland. Available online at:
http://www.math.ethz.ch/~wueth/Papers/2008_BF_V2.pdf
- [2] Arjas, E., *The Claims Reserving Problem in Non-Life Insurance: Some Structural Ideas* (1989), Astin Bulletin, vol. 19, no. 2, 139-152.
- [3] Björkwall, S., *Bootstrapping for claims reserve uncertainty in general insurance* (2009), Mathematical Statistics, Stockholm University, Research Report 2009:3, Licentiate thesis, ISSN 1650-0377. Available online at:
<http://www2.math.su.se/matstat/reports/seriea/2009/rep3/report.pdf>
- [4] Bornhuetter, R.L., and Ferguson, R.E., *The Actuary and IBNR* (1972), Proceedings of the Casualty Actuarial Society, 59(112), 181-195. Available online at:
<http://www.casact.org/pubs/proceed/proceed72/72181.pdf>
- [5] Christofides, S., *Regression Models Based on Log-incremental Payments* (1990), Claims Reserving Manual, vol.2, Institute of Actuaries, London.
- [6] Clarke, B. R., *Linear Models: the Theory and Application of Analysis of Variance* (2008), Wiley Series in Probability and Statistics, ISBN-10: 0470025662, 160-180.
- [7] Efron, B., Tibshirani, R., *An Introduction to the Bootstrap* (1994), Chapman & Hall/CRC, ISBN10: 0412042312, 31-199.
- [8] England, P. D., Verrall, R. J., *Analytic and bootstrap estimates of prediction errors in claims reserving* (1999), Insurance: Mathematics and Economics, vol. 25, 281-293.
- [9] England, P. D., Verrall, R. J., *Stochastic Claims Reserving in General Insurance* (2002), presented to the Institute of Actuaries. Available online at:
<http://www.emb.com/EMBDOTCOM/Netherlands/Nieuws%20-%20News/SCRinGI-EnglandVerrall.pdf>
- [10] England, P. D., Verrall, R. J., *More on Stochastic Reserving in General Insurance* (2004), Giro Convention, Killarney. Available online at:
www.actuaries.org.uk/_data/assets/powerpoint_doc/0006/22479/England.ppt
- [11] Hachemeister, C. A., Stanard, J. N., *IBNR Claims Count Estimation with Static Lag Functions* (1975), Spring Meeting of the Casualty Actuarial Society.
- [12] Li, J., *Comparison of Stochastic Reserving Methods* (2006), Australian Actuarial Journal, vol. 12, issue 4, 489-569.

- [13]Lowe, J., *A Practical Guide To Measuring Reserve Variability Using: Bootstrapping, Operational Time And A Distribution-Free Approach* (1994), General Insurance Convention, 157-196.
- [14]Mack, T., *A Simple Parametric Model for Rating Automobile Insurance or Estimating IBNR Claims Reserves* (1991), Astin Bulletin, vol. 22, no. 1, 93-109.
- [15]Mack, T., *Distribution-Free Calculation of the Standard Error of Chain Ladder Reserve Estimates* (1993), Astin Bulletin, vol. 23, no. 2.
- [16]Mack, T., *Measuring the Variability of Chain Ladder Reserve Estimates* (1993), CAS Prize Paper Competition. Available online at:
<http://www.casact.org/pubs/forum/94spforum/94spf101.pdf>
- [17]Mack, T., *The Standard Error of Chain Ladder Reserve Estimates: Recursive Calculation and Inclusion of Tail Factor* (1999), Astin Bulletin, vol. 29, no. 2, 361-366.
- [18]Mack, T., *Parameter Estimation of Bornhuetter/Ferguson* (2006:Fall), Casual Actuarial Society Forum. Available online at: <http://www.casact.org/pubs/forum/06fforum/145.pdf>
- [19]Mack, T., *The Prediction Error of Bornhuetter-Ferguson* (2008:Fall), Casualty Actuarial Society E-Forum, 222-240. Available online at:
<http://www.casact.org/pubs/forum/08fforum/10Mack%20.pdf>
- [20]Mack, T., Venter, G., *A Comparison of Stochastic Models that Reproduce Chain Ladder Reserve Estimates* (1999), Astin Colloquium, Tokyo, Japan – August 22-25, 263-273. Available online at: http://www.actuaries.org/ASTIN/Colloquia/Tokyo/Mack_Venter.pdf
- [21] Mack, T., Venter, G., *A Comparison of Stochastic Models that Reproduce Chain Ladder Reserve Estimates* (2000), Insurance: Mathematics and Economics, vol. 26, 101-107.
- [22]Pinheiro, P. J. R., Andrade e Silva, J. M., Centeno, Maria de L., *Bootstrap Methodology in Claims Reserving* (2001), Astin Colloquium International Actuarial Association – Brussels, Belgium. Available online at:
http://www.actuaries.org/ASTIN/Colloquia/Washington/Pinheiro_Silva_Centeno.pdf
- [23]Renshaw, A. E., *Chain-Ladder and Interactive Modelling (Claims Reserving and GLIM)* (1989), J. I. A., vol. 116, 559-587.
- [24]Stanard, J.N., *A Simulation Test of Prediction Errors of Loss Reserve Estimation Techniques* (1985), Proceedings of the Casualty Actuarial Society, 137&138, 124-148. Available online at:
<http://www.casact.org/pubs/proceed/proceed85/85124.pdf>
- [25]Schmidt, K. D., Zocher, M., *The Bornhuetter-Ferguson Principle* (2008), Casualty Actuarial Society, variance 2:1, 85-110.
- [26] Taylor, G. C., Ashe F. R., *Second Moments of Estimates of Outstanding Claims* (1983), Journal of Econometrics, vol. 23, 37-61.

- [27]Verrall, R. J., *Chain-Ladder and Maximum Likelihood* (1991), J. I. A., vol. 118, 489-499.
- [28]Verrall, R. J., *A Bayesian Generalised Linear Model for the Bornhuetter-Ferguson Method of Claims Reserving* (2001), Actuarial Research Paper, No. 139, ISBN10: 1901615626.
- [29]Verrall, R. J., *A Bayesian Generalised Linear Model for the Bornhuetter-Ferguson Method of Claims Reserving* (2005), North American Actuarial Journal, vol. 8, no. 3.
- [30]Wright, T. S., *A Stochastic Method for Claims Reserving in General Insurance* (1990), J.I.A., vol. 117, 667-731.
- [31]Wütrich, M.V., Bühlmann, H., and Furrer H., *Market-Consistent Actuarial Valuation* (2007), Springer, ISBN13: 9783540736424, 70-97.
- [32]Zehnwirth, B., *The Chain-Ladder Technique – a Stochastic Model* (1989), Claims Reserving Manual, vol. 2, Institute of Actuaries, London.

7. Appendix

7.1 Appendix I. The CL Prediction Error (Mack, 1993)

Lemma 1. Under (CL1) and (CL2) the estimators $\hat{f}_j$, $1 \leq k \leq n - 1$, are unbiased and uncorrelated.

Proof. Denote $B_k = \{C_{ij} | j \leq k, i + j \leq n + 1\}$, $1 \leq k \leq n$. Then

$$E(C_{i,k+1} | B_k) = E(C_{i,k+1} | C_{i1}, \dots, C_{ik}) = C_{ik} f_k$$

Also,

$$E(\hat{f}_k | B_k) = E\left(\frac{\sum_{i=1}^{n-k} C_{i,k+1}}{\sum_{i=1}^{n-k} C_{i,k}} \middle| B_k\right) = \frac{\sum_{i=1}^{n-k} E(C_{i,k+1} | B_k)}{\sum_{i=1}^{n-k} C_{i,k}} = \frac{\sum_{i=1}^{n-k} C_{i,k} f_k}{\sum_{i=1}^{n-k} C_{i,k}} = f_k$$

which gives

$$E(\hat{f}_k) = E(E(\hat{f}_k | B_k)) = E(f_k) = f_k, 1 \leq k \leq n - 1$$

The latter means that the estimator is unbiased. We can also write

$$E(\hat{f}_j \hat{f}_k) = E(E(\hat{f}_j \hat{f}_k | B_k)) = E(\hat{f}_j E(\hat{f}_k | B_k)) = E(\hat{f}_j) f_k = E(\hat{f}_j) E(\hat{f}_k)$$

which shows that the estimators are uncorrelated.

Q.E.D.

Theorem 1. Under the assumptions (CL1), (CL2), and (CL3) the mean square error $mse(\hat{R}_l)$ can be estimated by

$$\widehat{mse(\hat{R}_l)} = \hat{C}_{i,n}^2 \sum_{k=n+1-i}^{n-1} \frac{\hat{\sigma}_k^2}{\hat{f}_k^2} \left(\frac{1}{\hat{C}_{i,k}} + \frac{1}{\sum_{j=1}^{n-k} C_{j,k}} \right) \quad (1^*)$$

where $\widehat{C}_{i,k} = C_{i,n+1-i} \hat{f}_{n+1-i} \dots \hat{f}_{k-1}$, $k > n + 1 - i$ are the estimated values of the future $C_{i,k}$ and $\widehat{C}_{i,n+1-i} = C_{i,n+1-i}$.

Proof. The following abbreviations will be used

$$E_i(X) = E(X | C_{i,1}, \dots, C_{i,n+1-i})$$

$$Var_i(X) = Var(X | C_{i,1}, \dots, C_{i,n+1-i})$$

As shown in the text (see Formula (4)),

$$mse(\widehat{R}_i) = Var(C_{i,n}|D) + (E(C_{i,n}|D) - \widehat{C}_{i,n})^2 \quad (2^*)$$

The two summands will be estimated separately.

For the first term of (2*) repeated application of the CL assumptions (CL1) and (CL3), also noticing that $Var_i(C_{i,n+1-i}) = 0$, yields to

$$\begin{aligned} Var(C_{i,n}|D) &= Var_i(C_{i,n}) \\ &= E_i \left(Var(C_{i,n}|C_{i,1}, \dots, C_{i,n+1-i}) \right) + Var_i \left(E(C_{i,n}|C_{i,1}, \dots, C_{i,n+1-i}) \right) \\ &= E_i(C_{i,n-1}\sigma_{n-1}^2) + Var_i(C_{i,n-1}f_{n-1}) \\ &= E_i(C_{i,n-1})\sigma_{n-1}^2 + Var_i(C_{i,n-1})f_{n-1}^2 \\ &= E_i(C_{i,n-2})f_{n-2}\sigma_{n-1}^2 + E_i(C_{i,n-2})\sigma_{n-2}^2f_{n-1}^2 + Var_i(C_{i,n-2})f_{n-2}^2f_{n-1}^2 \\ &= \dots \\ &= C_{i,n+1-i} \sum_{k=n+1-i}^{n-1} f_{n+1-i} \dots f_{k-1} \sigma_k^2 f_{k+1}^2 \dots f_{n-1}^2 \end{aligned} \quad (3^*)$$

For the second term we have

$$(E(C_{i,n}|D) - \widehat{C}_{i,n})^2 = C_{i,n+1-i}^2 (f_{n+1-i} \dots f_{k-1} - \hat{f}_{n+1-i} \dots \hat{f}_{k-1})^2 \quad (4^*)$$

because $\widehat{C}_{i,n} = C_{i,n+1-i} \hat{f}_{n+1-i} \dots \hat{f}_{k-1}$ and

$$\begin{aligned} E(C_{i,n}|D) &= E_i(C_{i,n}) \\ &= E_i(E(C_{i,n}|C_{i,1}, \dots, C_{i,n-1})) \\ &= E_i(C_{i,n-1}f_{n-1}) = E_i(C_{i,n-1})f_{n-1} \\ &= \dots \\ &= E_i(C_{i,n+1-i})f_{n+1-i} \dots f_{n-1} \end{aligned}$$

Now, (3*) and (4*) have to be estimated. In the (3*) one can simply replace the unknown parameters f_k and σ_k^2 by their estimators $\hat{f}_k$ and $\hat{\sigma}_k^2$ and obtain

$$\begin{aligned}
C_{i,n+1-i} & \sum_{k=n+1-i}^{n-1} \hat{f}_{n+1-i} \cdots \hat{f}_{k-1} \hat{\sigma}_k^2 \hat{f}_{k+1}^2 \cdots \hat{f}_{n-1}^2 \\
& = C_{i,n+1-i} (\hat{\sigma}_{n+1-i}^2 \hat{f}_{n+2-i}^2 \cdots \hat{f}_{n-1}^2 + \hat{f}_{n+1-i} \hat{\sigma}_{n+2-i}^2 \hat{f}_{n+3-i}^2 \cdots \hat{f}_{n-1}^2 + \cdots + \hat{f}_{n+1-i} \cdots \hat{f}_{n-2} \hat{\sigma}_{n-1}^2) \\
& = C_{i,n+1-i}^2 \hat{f}_{n+1-i}^2 \cdots \hat{f}_{n-1}^2 \left(\frac{\frac{\hat{\sigma}_{n+1-i}^2}{\hat{f}_{n+1-i}^2}}{\hat{C}_{i,n+1-i}} + \frac{\frac{\hat{\sigma}_{n+2-i}^2}{\hat{f}_{n+2-i}^2}}{\hat{C}_{i,n+1-i} \hat{f}_{n+1-i}^2} + \cdots + \frac{\frac{\hat{\sigma}_{n-1}^2}{\hat{f}_{n-1}^2}}{\hat{C}_{i,n+1-i} \hat{f}_{n+1-i}^2 \cdots \hat{f}_{n-2}^2} \right) \\
& = \hat{C}_{i,n}^2 \sum_{k=n+1-i}^{n-1} \frac{\hat{\sigma}_k^2 / \hat{f}_k^2}{\hat{C}_{i,k}}
\end{aligned}$$

The same cannot be done for the (4*), since this kind of replacement would just give 0. Therefore, a different approach must be used. It can be written

$$F = f_{n+1-i} \cdots f_{n-1} - \hat{f}_{n+1-i} \cdots \hat{f}_{n-1} = S_{n+1-i} + \cdots + S_{n-1}$$

with $S_k = \hat{f}_{n+1-i} \cdots \hat{f}_{k-1} (f_k - \hat{f}_k) f_{k+1} \cdots f_{n-1}$ and therefore

$$F^2 = (S_{n+1-i} + \cdots + S_{n-1})^2 = \sum_{k=n+1-i}^{n-1} S_k^2 + 2 \sum_{j < k} S_j S_k$$

Now we replace S_k^2 with $E(S_k^2 | B_k)$ and $S_j S_k, j < k$, with $E(S_j S_k | B_k)$, where $B_k = \{C_{ij} | j \leq k, i + j \leq n + 1\}, 1 \leq k \leq n$. This means that we approximate S_k^2 and $S_j S_k$ by averaging over as little data as possible such that as many values C_{ik} as possible from the observed data are kept fixed. Since $E(f_k - \hat{f}_k | B_k) = 0$, it follows that $E(S_j S_k | B_k) = 0$ for $j < k$. Since

$$\begin{aligned}
E((f_k - \hat{f}_k)^2 | B_k) & = \text{Var}(\hat{f}_k | B_k) \\
& = \text{Var}\left(\frac{\sum_{j=1}^{n-k} C_{j,k+1}}{\sum_{j=1}^{n-k} C_{j,k}} \middle| B_k\right) \\
& = \sum_{j=1}^{n-k} \frac{\text{Var}(C_{j,k+1} | B_k)}{(\sum_{j=1}^{n-k} C_{j,k})^2} \\
& = \frac{\sigma_k^2 \sum_{j=1}^{n-k} C_{j,k}}{(\sum_{j=1}^{n-k} C_{j,k})^2} \\
& = \frac{\sigma_k^2}{\sum_{j=1}^{n-k} C_{j,k}}
\end{aligned}$$

it follows that

$$E(S_k^2|B_k) = (\hat{f}_{n+1-i}^2 \cdots \hat{f}_{k-1}^2 \sigma_k^2 f_{k+1}^2 \cdots f_{n-1}^2) \frac{1}{\sum_{j=1}^{n-k} C_{j,k}}$$

From above, $F^2 = (\sum S_k)^2$, and we replace it with $\sum_k E(S_k^2|B_k)$ and because all terms of this sum are positive we now can replace all unknown parameters f_k and σ_k^2 by their unbiased estimators $\hat{f}_k$ and $\hat{\sigma}_k^2$. Altogether, we estimate $F^2 = (f_{n+1-i} \cdots f_{k-1} - \hat{f}_{n+1-i} \cdots \hat{f}_{k-1})^2$ by

$$\begin{aligned} & \sum_{k=n+1-i}^{n-1} ((\hat{f}_{n+1-i}^2 \cdots \hat{f}_{k-1}^2 \hat{\sigma}_k^2 \hat{f}_{k+1}^2 \cdots \hat{f}_{n-1}^2) \frac{1}{\sum_{j=1}^{n-k} C_{j,k}}) \\ &= \left(\hat{\sigma}_{n+1-i}^2 \hat{f}_{n+2-i}^2 \cdots \hat{f}_{n-1}^2 \frac{1}{\sum_{j=1}^{n-(n+1-i)} C_{j,n+1-i}} \right. \\ & \quad + \hat{f}_{n+1-i} \hat{\sigma}_{n+2-i}^2 \hat{f}_{n+3-i}^2 \cdots \hat{f}_{n-1}^2 \frac{1}{\sum_{j=1}^{n-(n+2-i)} C_{j,n+2-i}} + \cdots \\ & \quad \left. + \hat{f}_{n+1-i} \cdots \hat{f}_{n-2} \hat{\sigma}_{n-1}^2 \frac{1}{\sum_{j=1}^{n-(n-1)} C_{j,n-1}} \right) \\ &= \hat{f}_{n+1-i}^2 \cdots \hat{f}_{n-1}^2 \left(\frac{\frac{\hat{\sigma}_{n+1-i}^2}{\hat{f}_{n+1-i}^2}}{\sum_{j=1}^{n-(n+1-i)} C_{j,n+1-i}} + \frac{\frac{\hat{\sigma}_{n+2-i}^2}{\hat{f}_{n+2-i}^2}}{\sum_{j=1}^{n-(n+2-i)} C_{j,n+2-i}} + \cdots + \frac{\frac{\hat{\sigma}_{n-1}^2}{\hat{f}_{n-1}^2}}{\sum_{j=1}^{n-(n-1)} C_{j,n-1}} \right) \\ &= \hat{f}_{n+1-i}^2 \cdots \hat{f}_{n-1}^2 \sum_{k=n+1-i}^{n-1} \frac{\hat{\sigma}_k^2 / \hat{f}_k^2}{\sum_{j=1}^{n-k} C_{j,k}} \end{aligned}$$

Putting the two estimated summands together leads to the estimator (1*) stated in the theorem:

$$\begin{aligned} \widehat{mse}(\hat{R}_i) &= \hat{C}_{i,n}^2 \sum_{k=n+1-i}^{n-1} \frac{\hat{\sigma}_k^2 / \hat{f}_k^2}{\hat{C}_{i,k}} + C_{i,n+1-i}^2 \hat{f}_{n+1-i}^2 \cdots \hat{f}_{n-1}^2 \sum_{k=n+1-i}^{n-1} \frac{\hat{\sigma}_k^2 / \hat{f}_k^2}{\sum_{j=1}^{n-k} C_{j,k}} \\ &= \hat{C}_{i,n}^2 \sum_{k=n+1-i}^{n-1} \frac{\hat{\sigma}_k^2}{\hat{f}_k^2} \left(\frac{1}{\hat{C}_{i,k}} + \frac{1}{\sum_{j=1}^{n-k} C_{j,k}} \right) \end{aligned}$$

Q.E.D.

Corollary. With the assumptions and notations of Theorem 1 the mean squared error of the overall reserve estimate $\hat{R} = \hat{R}_2 + \cdots + \hat{R}_n$ can be estimated by

$$\widehat{mse}(\hat{R}) = \sum_{i=2}^n \left\{ (s.e.(\hat{R}_i))^2 + \hat{C}_{i,n} (\sum_{j=i+1}^n \hat{C}_{j,n}) \sum_{k=n+1-i}^{n-1} \frac{\frac{2\hat{\sigma}_k^2}{\hat{f}_k^2}}{\sum_{l=1}^{n-k} C_{l,k}} \right\} \quad (5^*)$$

Proof. It can be written by definition

$$\begin{aligned}
 mse\left(\sum_{i=2}^n \hat{R}_i\right) &= E\left(\left(\sum_{i=2}^n \hat{R}_i - \sum_{i=2}^n R_i\right)^2 \middle| D\right) \\
 &= E\left(\left(\sum_{i=2}^n \hat{C}_{i,n} - \sum_{i=2}^n C_{i,n}\right)^2 \middle| D\right) \\
 &= Var\left(\sum_{i=2}^n C_{i,n} \middle| D\right) + \left(E\left(\sum_{i=2}^n C_{i,n} \middle| D\right) - \sum_{i=2}^n \hat{C}_{i,n}\right)^2
 \end{aligned}$$

The independence of accident quarters (CL2) gives the following

$$Var\left(\sum_{i=2}^n C_{i,n} \middle| D\right) = \sum_{i=2}^n Var(C_{i,n}|D)$$

And $Var(C_{i,n}|D) = C_{i,n+1-i} \sum_{k=n+1-i}^{n-1} f_{n+1-i} \dots f_{k-1} \sigma_k^2 f_{k+1}^2 \dots f_{n-1}^2$ from above. Also,

$$\begin{aligned}
 \left(E\left(\sum_{i=2}^n C_{i,n} \middle| D\right) - \sum_{i=2}^n \hat{C}_{i,n}\right)^2 &= \left(\sum_{i=2}^n (E(C_{i,n}|D) - \hat{C}_{i,n})\right)^2 \\
 &= \sum_{i,j} (E(C_{i,n}|D) - \hat{C}_{i,n})(E(C_{j,n}|D) - \hat{C}_{j,n}) \\
 &= \sum_{i,j} C_{i,n+1-i} C_{j,n+1-j} F_i F_j
 \end{aligned}$$

with $F_i = f_{n+1-i} \dots f_{n-1} - \hat{f}_{n+1-i} \dots \hat{f}_{n-1}$.

Remember from Theorem 1 formulae (2*) and (4*) that $mse(\hat{R}_i) = Var(C_{i,n}|D) + (C_{i,n+1-i} F_i)^2$.

Then it can be written:

$$mse\left(\sum_{i=2}^n \hat{R}_i\right) = \sum_{i=2}^n mse(\hat{R}_i) + \sum_{2 \leq i < j \leq n} 2 C_{i,n+1-i} C_{j,n+1-j} F_i F_j$$

Let us use the same procedure as for F^2 in the proof to Theorem 1 here again for $F_i F_j$. Assume $i < j$ and $j - i = diff$. Then we write

$$\begin{aligned}
F_i F_j &= (S_{n+1-i} + \dots + S_{n-1})(S_{n+1-j} + \dots + S_{n-1}) \\
&= (S_{n+1-i} + \dots + S_{n-1})(S_{n+1-j+diff} + \dots + S_{n-1}) \\
&\quad + (S_{n+1-i} + \dots + S_{n-1})(S_{n+1-j} + \dots + S_{n+1-j+diff-1})
\end{aligned}$$

Notice that the first multiplication does not give 0 on average since $n+1-i = n+1-j+diff$. But the second multiplication is equal to 0 on average because multiplying terms are always different.

$$S_{n+1-i} = (f_{n+1-i} - \hat{f}_{n+1-i})f_{n+2-i} \dots f_{n-1}$$

$$S_{n+1-j+diff} = \hat{f}_{n+1-j} \dots \hat{f}_{n+1-j+diff-1}(f_{n+1-j+diff} - \hat{f}_{n+1-j+diff})f_{n+2-j+diff} \dots f_{n-1}$$

Therefore,

$$F_i F_j = \sum_{k=n+1-i}^{n-1} \hat{f}_{n+1-j} \dots \hat{f}_{n-i} S_k^2 = \hat{f}_{n+1-j} \dots \hat{f}_{n-i} \sum_{k=n+1-i}^{n-1} S_k^2$$

And $\sum_{k=n+1-i}^{n-1} S_k^2 = \hat{f}_{n+1-i}^2 \dots \hat{f}_{n-1}^2 \sum_{k=n+1-i}^{n-1} \frac{\hat{\sigma}_k^2 / \hat{f}_k^2}{\sum_{j=1}^{n-k} C_{j,k}}$, as computed earlier (see proof of Theorem 1).

Finally, we can estimate $F_i F_j$ by

$$\hat{f}_{n+1-j} \dots \hat{f}_{n-i} \hat{f}_{n+1-i}^2 \dots \hat{f}_{n-1}^2 \sum_{k=n+1-i}^{n-1} \frac{\hat{\sigma}_k^2 / \hat{f}_k^2}{\sum_{j=1}^{n-k} C_{j,k}}$$

Inserting the latter expression into $mse(\sum_{i=2}^n \hat{R}_i)$ formula the stated expression (5*) is obtained:

$$\begin{aligned}
mse\left(\sum_{i=2}^n \hat{R}_i\right) &= \sum_{i=2}^n mse(\hat{R}_i) + 2 \sum_{i=2}^n \sum_{j=i+1}^n C_{i,n+1-i} C_{j,n+1-j} \hat{f}_{n+1-j} \dots \hat{f}_{n-i} \hat{f}_{n+1-i}^2 \dots \hat{f}_{n-1}^2 \sum_{k=n+1-i}^{n-1} \frac{\hat{\sigma}_k^2 / \hat{f}_k^2}{\sum_{j=1}^{n-k} C_{j,k}} \\
&= \sum_{i=2}^n \left\{ mse(\hat{R}_i) + C_{i,n+1-i} \sum_{j=i+1}^n C_{j,n+1-j} \hat{f}_{n+1-j} \dots \hat{f}_{n-i} \hat{f}_{n+1-i}^2 \dots \hat{f}_{n-1}^2 \sum_{k=n+1-i}^{n-1} \frac{\hat{\sigma}_k^2 / \hat{f}_k^2}{\sum_{j=1}^{n-k} C_{j,k}} \right\} \\
&= \sum_{i=2}^n \left\{ (s.e.(\hat{R}_i))^2 + \hat{C}_{i,n} \sum_{j=i+1}^n \hat{C}_{j,n} \sum_{k=n+1-i}^{n-1} \frac{\hat{\sigma}_k^2 / \hat{f}_k^2}{\sum_{j=1}^{n-k} C_{j,k}} \right\}
\end{aligned}$$

Q.E.D.

7.2 Appendix II. The BF Prediction Error (Mack, 2008)

Theorem 2. Under the assumptions (BF1), (BF2), and (BF3) the mean square error $mse(\hat{R}_i)$ can be estimated by

$$mse(\hat{R}_i^{BF}) = \hat{U}_i(\hat{s}_{n+2-i}^2 + \dots + \hat{s}_{n+1}^2) + \left(\hat{U}_i^2 + (s.e.(\hat{U}_i))^2 \right) (s.e.(\hat{z}_{n+1-i}^*))^2 + (s.e.(\hat{U}_i))^2 (1 - \hat{z}_{n+1-i}^*)^2 \quad (1^{**})$$

Proof. Notice, the mean square error of prediction of the BF reserve estimate is the sum of the squared estimation error $Var(\hat{R}_i^{BF})$ and of the squared process error $Var(R_i)$. The two terms will be estimated separately.

First, the process error:

$$\begin{aligned} Var(R_i) &= Var(S_{i,n+2-i} + \dots + S_{i,n+1}) = Var(S_{i,n+2-i}) + \dots + Var(S_{i,n+1}) \\ &= x_i(s_{n+2-i}^2 + \dots + s_{n+1}^2) \end{aligned}$$

This will be estimated by

$$\widehat{Var}(R_i) = \hat{U}_i(\hat{s}_{n+2-i}^2 + \dots + \hat{s}_{n+1}^2). \quad (2^{**})$$

Now, we will turn to the second, more difficult term – the estimation error $Var(\hat{R}_i^{BF})$. First, the following general formula will be used:

$$Var(XY) = (E(X))^2 Var(Y) + Var(X)Var(Y) + Var(X)(E(Y))^2$$

In our case it becomes:

$$\begin{aligned} Var(\hat{R}_i^{BF}) &= Var(\hat{U}_i(1 - \hat{z}_{n+1-i}^*)) = \\ &= (E(\hat{U}_i))^2 Var(\hat{z}_{n+1-i}^*) + Var(\hat{U}_i)Var(\hat{z}_{n+1-i}^*) + Var(\hat{U}_i)(1 - E(\hat{z}_{n+1-i}^*))^2 = \\ &= (x_i^2 + Var(\hat{U}_i)) Var(\hat{z}_{n+1-i}^*) + Var(\hat{U}_i)(1 - z_{n+1-i})^2 \quad (3^{**}) \end{aligned}$$

In the latter formula the estimators for x_i and z_{n+1-i} have already been derived, but estimates for $Var(\hat{U}_i)$ and $Var(\hat{z}_{n+1-i}^*)$ are still not found. Therefore, the next step will be to find the best estimates for the precision of $\hat{U}_i$ and $\hat{z}_{n+1-i}^*$.

Mack (2008) suggests that like $\hat{U}_i$ itself, $s.e.(\hat{U}_i)$ would be best if obtained from repricing of the business. But one should be careful here. The usual standard deviation formula

$$(s.e.(\hat{U}_i))^2 = \frac{v_i}{n-1} \sum_{j=1}^n v_j \left(\frac{\hat{U}_i}{v_j} - \hat{q} \right)^2, \text{ with } \hat{q} = \frac{\sum_{j=1}^n \hat{U}_j}{\sum_{j=1}^n v_j}$$

can be used only if the initial estimates $\hat{U}_i$ can be assumed to be uncorrelated. The market cycle of premium adequacy should be removed where possible, otherwise we would overestimate $s.e.(\hat{U}_i)$. To remove the market cycle the on-level premiums $\tilde{v}_j$ should be used. In the case when positive correlation exists between the $\hat{U}_i$'s the term $n-1$ should be replaced by $n-\sqrt{n}$ for a constant correlation coefficient $\hat{\rho}_{ij}^U = \frac{1}{\sqrt{n}}$, or by $n-\sqrt{2n}$ for a decreasing correlation coefficient $\hat{\rho}_{ij}^U = \frac{1}{1+|i-j|}$, or by a precise formula $n - \sum_{i,j} \rho_{ij}^U \sqrt{\frac{v_i v_j}{v_+ v_+}}$, with $v_+ = \sum_{i=1}^n v_i$. After the estimation, one should examine the plausibility of the estimates. One way to examine it is to assume normal distribution and see if the interval $(\hat{U}_i - 2 s.e.(\hat{U}_i), \hat{U}_i + 2 s.e.(\hat{U}_i))$ will contain the true $E(\hat{U}_i)$ with 95% probability.

In order to estimate the precision of $\hat{z}_{n+1-i}^*$, we first write out what it is:

$$Var(1 - \hat{z}_{n+1-i}^*) = Var(\hat{z}_{n+1-i}^*) = Var(\hat{y}_1^* + \dots + \hat{y}_{n+1-i}^*) = Var(\hat{y}_{n+2-i}^* + \dots + \hat{y}_{n+1}^*)$$

It does not contradict to replace $Var(\hat{y}_1^* + \dots + \hat{y}_{n+1-i}^*)$ with $Var(\hat{y}_1^*) + \dots + Var(\hat{y}_{n+1-i}^*)$, but the latter sum increases with each term, which is not the case since $Var(\hat{y}_1^* + \dots + \hat{y}_{n+1}^*) = Var(1) = 0$. The solution will be to take $Var(\hat{y}_1^*) + \dots + Var(\hat{y}_k^*)$ for small k and $Var(\hat{y}_{k+1}^*) + \dots + Var(\hat{y}_{n+1}^*)$ for large k . More precisely,

$$Var(\hat{z}_k^*) = \min(Var(\hat{y}_1^*) + \dots + Var(\hat{y}_k^*), Var(\hat{y}_{k+1}^*) + \dots + Var(\hat{y}_{n+1}^*))$$

Due to the previous estimation procedure of the parameters, we can proceed as follows:

$$Var(\hat{y}_k^*) \approx Var\left(\frac{\sum_{j=1}^{n+1-i} S_{j,k}}{\sum_{j=1}^{n+1-i} x_j}\right) = \frac{s_k^2}{\sum_{j=1}^{n+1-i} x_j}, 1 \leq k \leq n$$

This allows us to estimate $Var(\hat{y}_k^*)$ by

$$(s.e.(\hat{y}_k^*))^2 = \frac{\hat{s}_k^{2*}}{\sum_{j=1}^{n+1-i} \hat{U}_j}, 1 \leq k \leq n$$

Notice, that the value of $s.e.(\hat{y}_k^*)$ must come from outside. Otherwise, this way of estimating will not lead to reliable estimates. If no outside value is available, a plausible choice is to assume a normal distribution with 95% probability within the interval $(0; 2\hat{y}_{n+1}^*)$, and coefficient of variation being 50%. Then $s.e.(\hat{y}_k^*) = 0.5\hat{y}_{n+1}^*$. Then

$$(s.e.(\hat{z}_k^*))^2 = \min \left((s.e.(\hat{y}_1^*))^2 + \dots + (s.e.(\hat{y}_k^*))^2, (s.e.(\hat{y}_{k+1}^*))^2 + \dots + (s.e.(\hat{y}_{n+1}^*))^2 \right)$$

Here again the plausibility of the estimate should be checked, for example, in the same way as it was done for $s.e.(\hat{U}_i)$.

Now, putting both process (2**) and estimation error (3**) together the mean square error of prediction (1**) is obtained

$$\begin{aligned} mse(\hat{R}_i^{BF}) &= \hat{U}_i(\hat{s}_{n+2-i}^2 + \dots + \hat{s}_{n+1}^2) + \left(\hat{U}_i^2 + (s.e.(\hat{U}_i))^2 \right) (s.e.(\hat{z}_{n+1-i}^*))^2 \\ &\quad + (s.e.(\hat{U}_i))^2 (1 - \hat{z}_{n+1-i}^*)^2 \end{aligned}$$

Q.E.D.

Corollary. With the assumptions and notations of Theorem 2 the mean squared error of the overall reserve estimate $\hat{R}^{BF} = \hat{R}_1^{BF} + \dots + \hat{R}_n^{BF}$ can be estimated by

$$mse(\hat{R}^{BF}) = \sum_{i=1}^n \hat{U}_i(\hat{s}_{n+2-i}^2 + \dots + \hat{s}_{n+1}^2) + \sum_{i=1}^n Var(\hat{R}_i^{BF}) + 2 \sum_{i < j} Cov(\hat{R}_i^{BF}, \hat{R}_j^{BF}) \quad (4^{**})$$

Proof. The mean square error of prediction of the overall reserve estimate has again two terms, the process error and the estimation error. Since, according to (BF1), all the increments are independent, the process error is

$$Var(R) = Var(R_1) + \dots + Var(R_n) = \sum_{i=1}^n \hat{U}_i(\hat{s}_{n+2-i}^2 + \dots + \hat{s}_{n+1}^2)$$

The estimation error involves more complicated calculations because $\hat{R}_1^{BF}, \dots, \hat{R}_n^{BF}$ are positively correlated via the common parameter estimates $\hat{y}_k^*$.

$$Var(\hat{R}^{BF}) = \sum_{i=1}^n Var(\hat{R}_i^{BF}) + 2 \sum_{i < j} Cov(\hat{R}_i^{BF}, \hat{R}_j^{BF})$$

The general formula to compute covariances for independent sets $\{X, W\}$ and $\{Y, Z\}$ will be used

$$Cov(XY, WZ) = Cov(X, W)E(Y)E(Z) + Cov(X, W)Cov(Y, Z) + E(X)E(W)Cov(Y, Z)$$

Then we have

$$\begin{aligned}
Cov(\hat{R}_i^{BF}, \hat{R}_j^{BF}) &= Cov(\hat{U}_i(1 - \hat{z}_{n+1-i}^*), \hat{U}_j(1 - \hat{z}_{n+1-j}^*)) \\
&= Cov(\hat{U}_i, \hat{U}_j)E(1 - \hat{z}_{n+1-i}^*)E(1 - \hat{z}_{n+1-j}^*) \\
&\quad + Cov(\hat{U}_i, \hat{U}_j)Cov(1 - \hat{z}_{n+1-i}^*, 1 - \hat{z}_{n+1-j}^*) \\
&\quad + E(\hat{U}_i)E(\hat{U}_j)Cov(1 - \hat{z}_{n+1-i}^*, 1 - \hat{z}_{n+1-j}^*)
\end{aligned}$$

The term $Cov(\hat{U}_i, \hat{U}_j)Cov(1 - \hat{z}_{n+1-i}^*, 1 - \hat{z}_{n+1-j}^*)$ is of lower order and can be omitted. Therefore we are left with

$$\begin{aligned}
Cov(\hat{R}_i^{BF}, \hat{R}_j^{BF}) &= Cov(\hat{U}_i, \hat{U}_j)E(1 - \hat{z}_{n+1-i}^*)E(1 - \hat{z}_{n+1-j}^*) \\
&\quad + E(\hat{U}_i)E(\hat{U}_j)Cov(1 - \hat{z}_{n+1-i}^*, 1 - \hat{z}_{n+1-j}^*) \\
&= \rho_{ij}^U \sqrt{Var(\hat{U}_i)Var(\hat{U}_j)}E(1 - \hat{z}_{n+1-i}^*)E(1 - \hat{z}_{n+1-j}^*) \\
&\quad + \rho_{ij}^Z \sqrt{Var(\hat{z}_{n+1-i}^*)Var(\hat{z}_{n+1-j}^*)}E(\hat{U}_i)E(\hat{U}_j)
\end{aligned}$$

with the correlation coefficients

$$\begin{aligned}
\rho_{ij}^U &= \frac{Cov(\hat{U}_i, \hat{U}_j)}{\sqrt{Var(\hat{U}_i)Var(\hat{U}_j)}} \\
\rho_{ij}^Z &= \frac{Cov(1 - \hat{z}_{n+1-i}^*, 1 - \hat{z}_{n+1-j}^*)}{\sqrt{Var(\hat{z}_{n+1-i}^*)Var(\hat{z}_{n+1-j}^*)}}
\end{aligned}$$

These correlation coefficients are what we have left to compute. One can estimate them from the data, but if this way is not possible the $\hat{\rho}_{ij}^U$ can be obtained as in the explained procedure earlier and

$$\hat{\rho}_{ij}^Z = \sqrt{\frac{\hat{z}_{n+1-j}^*(1 - \hat{z}_{n+1-i}^*)}{\hat{z}_{n+1-i}^*(1 - \hat{z}_{n+1-j}^*)}}$$

The latter approach will be used later in this work analyzing the personal accident claims. Now, the final result can be stated – the prediction error for the total BF reserve estimate (4**):

$$mse(\hat{R}^{BF}) = \sum_{i=1}^n \hat{U}_i(\hat{s}_{n+2-i}^{2*} + \dots + \hat{s}_{n+1}^{2*}) + \sum_{i=1}^n Var(\hat{R}_i^{BF}) + 2 \sum_{i < j} Cov(\hat{R}_i^{BF}, \hat{R}_j^{BF})$$

with

$$\widehat{Cov}(\hat{R}_i^{BF}, \hat{R}_j^{BF}) = \hat{\rho}_{ij}^U s.e.(\hat{U}_i) s.e.(\hat{U}_j)(1 - \hat{z}_{n+1-i}^*)(1 - \hat{z}_{n+1-j}^*) + \hat{\rho}_{ij}^Z s.e.(\hat{z}_{n+1-i}^*) s.e.(\hat{z}_{n+1-j}^*) \hat{U}_i \hat{U}_j$$

7.3 Appendix III. (CL1) and (CL3) check

7.3.1 Appendix III.a – Incurred cumulative claims linearity (CL1) check

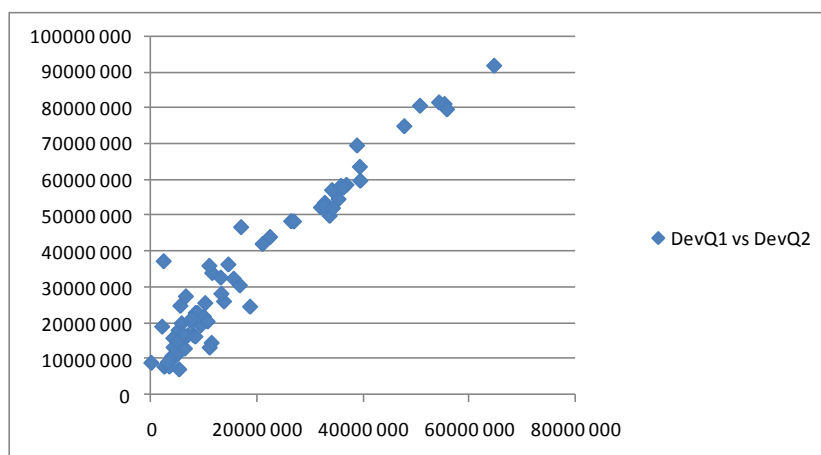


Figure 1*. Incurred claims at the 1st development quarter against the ones at the 2nd development quarter.

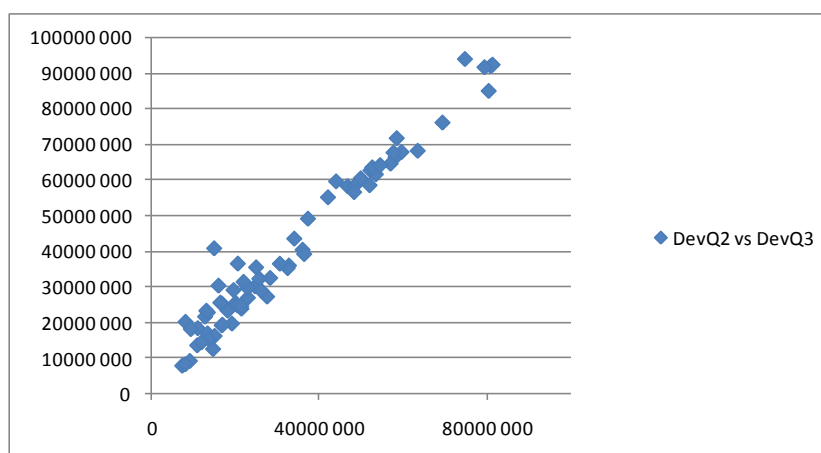


Figure 2*. Incurred claims at the 2nd development quarter against the ones at the 3rd development quarter.

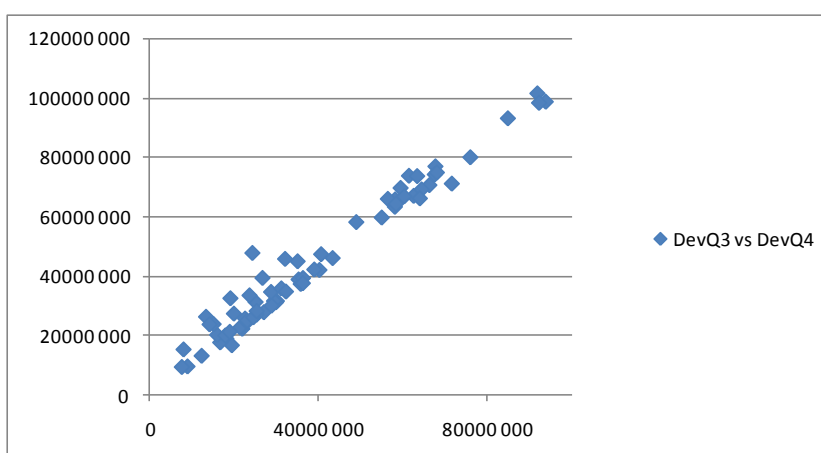


Figure 3*. Incurred claims at the 3rd development quarter against the ones at the 4th development quarter.

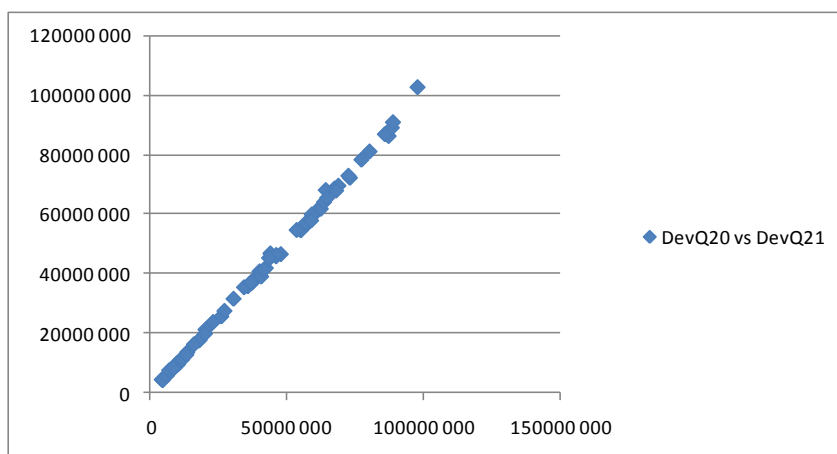


Figure 4*. Incurred claims at the 20th development quarter against the ones at the 21st development quarter.

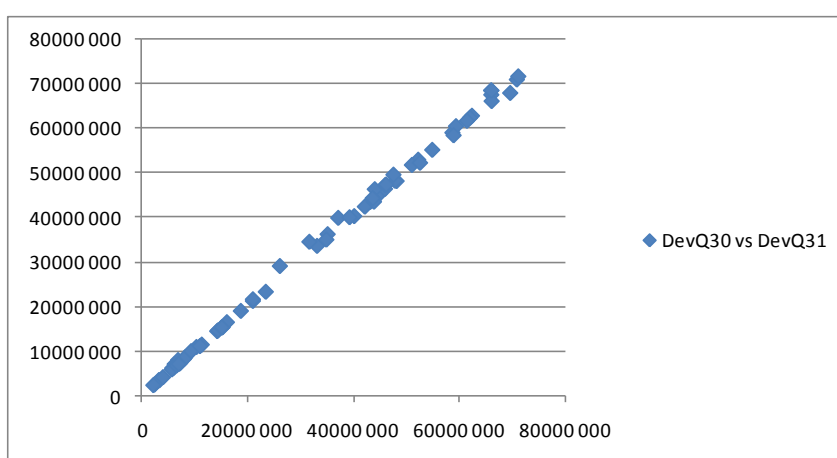


Figure 5*. Incurred claims at the 30th development quarter against the ones at the 31st development quarter.

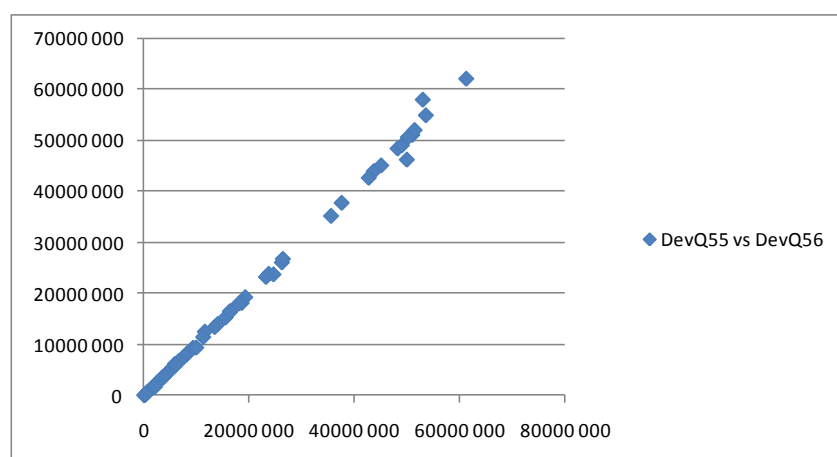


Figure 6*. Incurred claims at the 55th development quarter against the ones at the 56th development quarter.

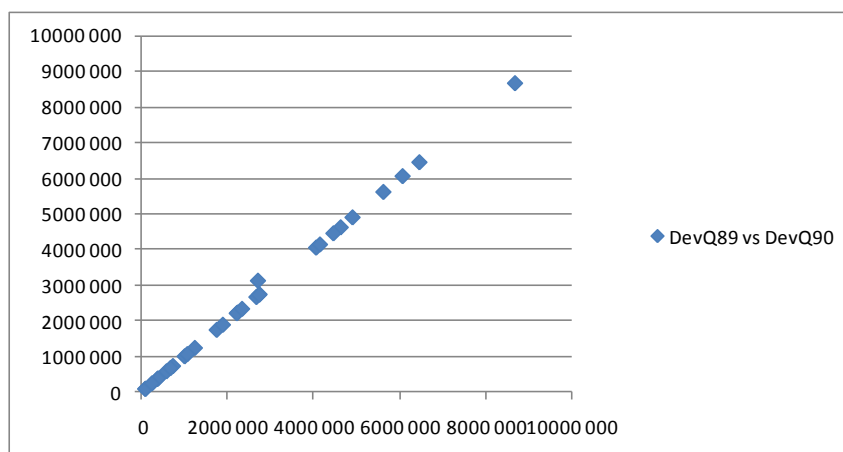


Figure 7*. Incurred claims at the 89th development quarter against the ones at the 90th development quarter.

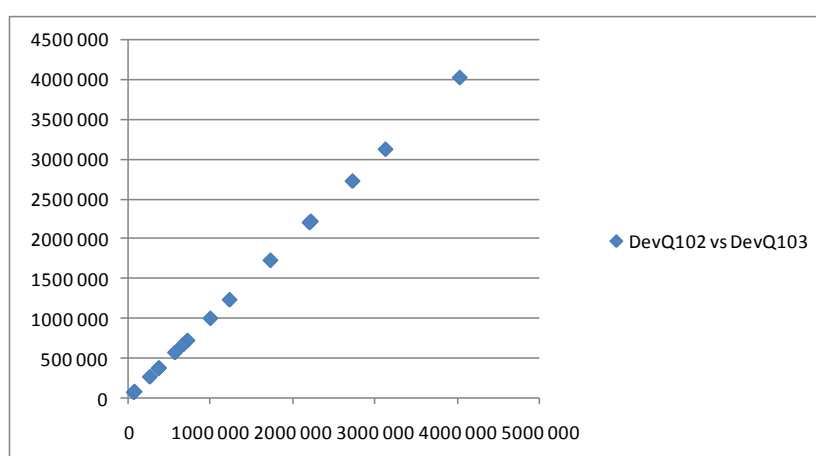


Figure 8*. Incurred claims at the 102nd development quarter against the ones at the 103rd development quarter.

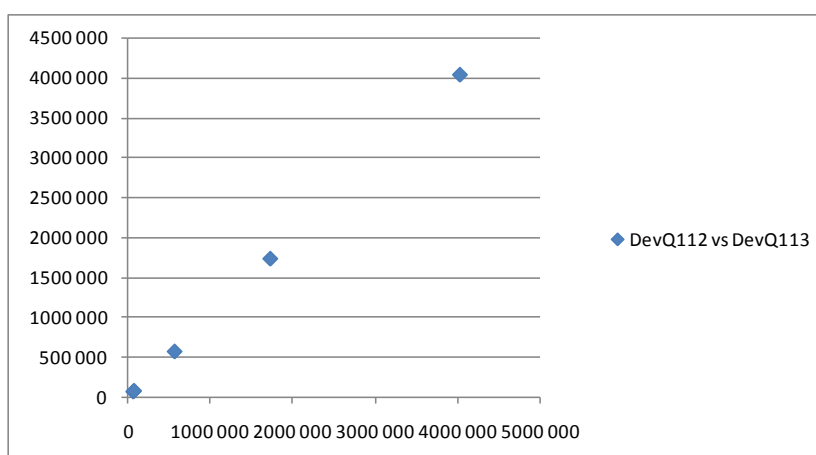


Figure 9*. Incurred claims at the 112th development quarter against the ones at the 113th development quarter.

7.3.2 Appendix III.b – Paid cumulative claims linearity (CL1) check

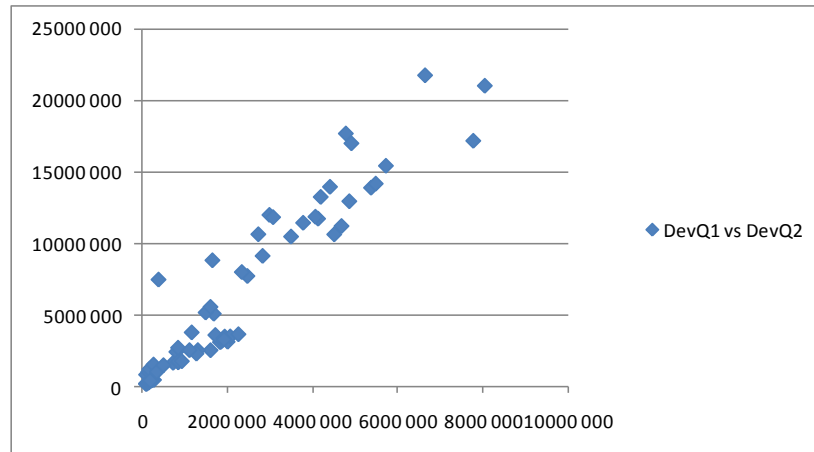


Figure 10*. Paid claims at the 1st development quarter against the ones at the 2nd development quarter.

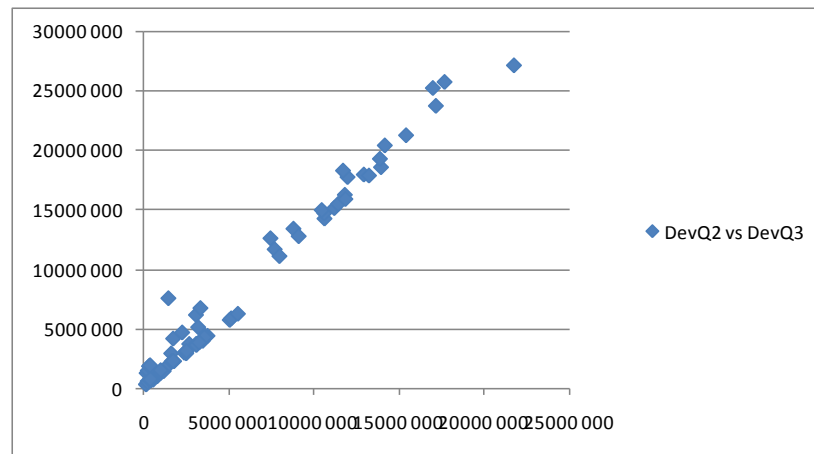


Figure 11*. Paid claims at the 2nd development quarter against the ones at the 3rd development quarter.

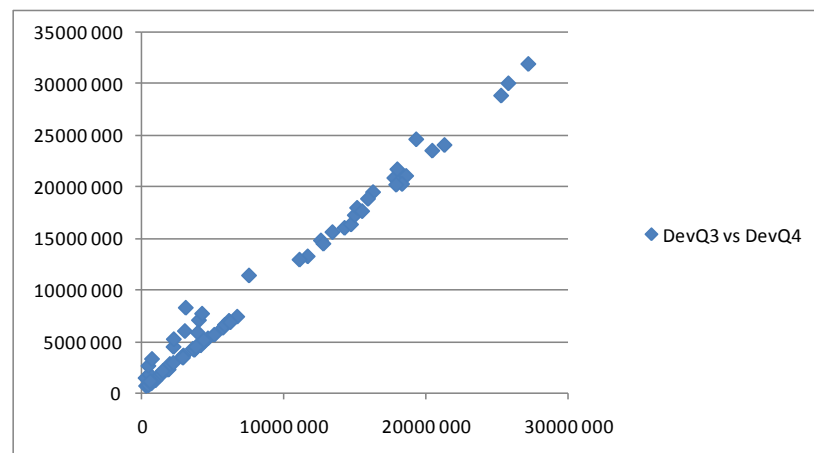


Figure 12*. Paid claims at the 3rd development quarter against the ones at the 4th development quarter.

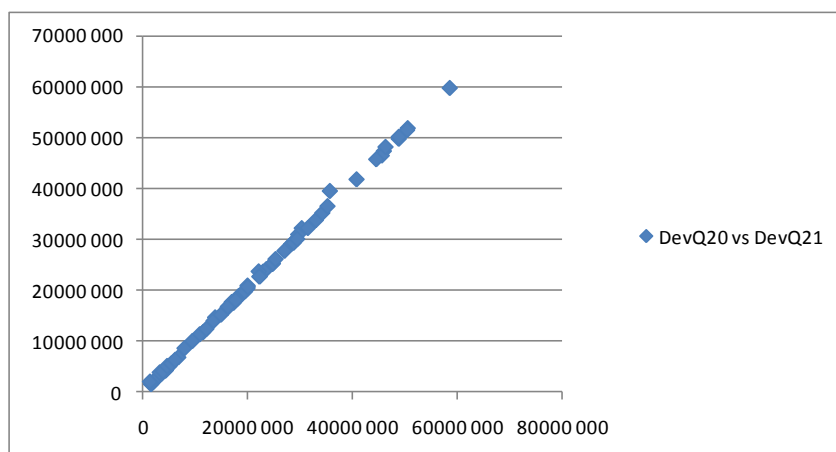


Figure 13*. Paid claims at the 20th development quarter against the ones at the 21st development quarter.

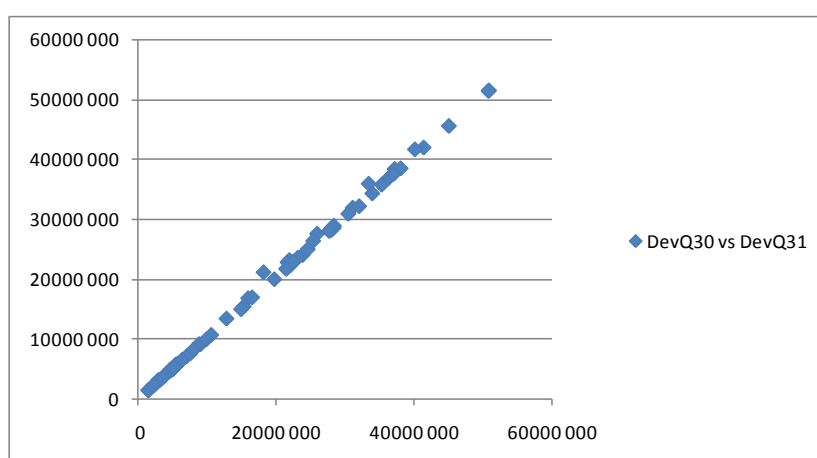


Figure 14*. Paid claims at the 30th development quarter against the ones at the 31st development quarter.

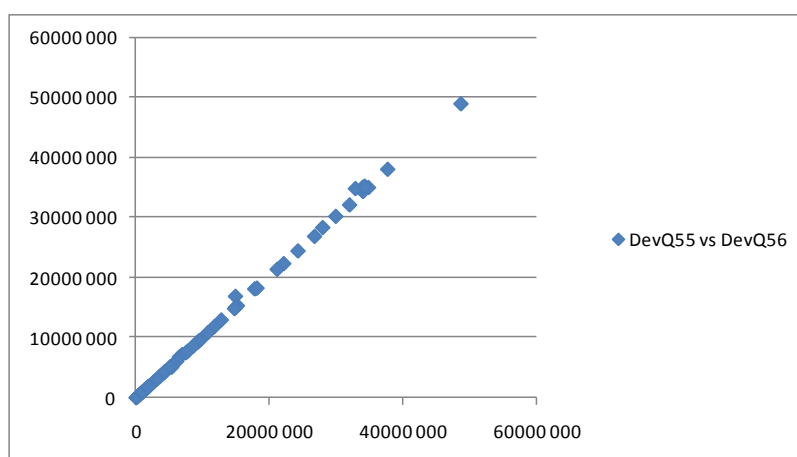


Figure 15*. Paid claims at the 55th development quarter against the ones at the 56th development quarter.

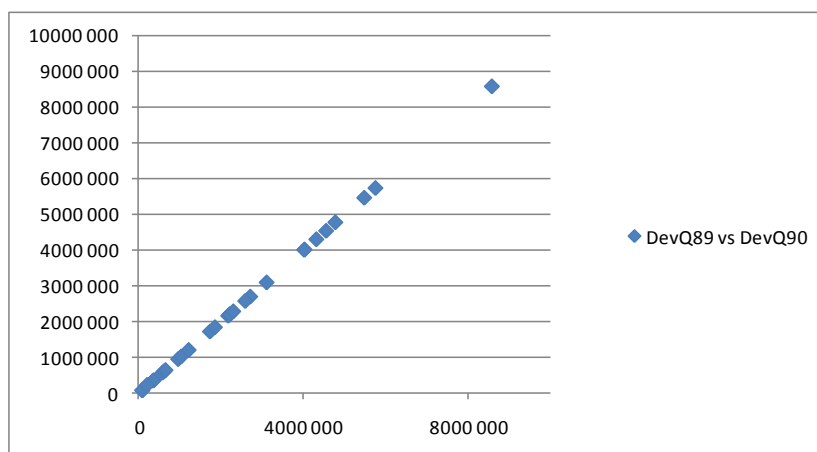


Figure 16*. Paid claims at the 89th development quarter against the ones at the 90th development quarter.

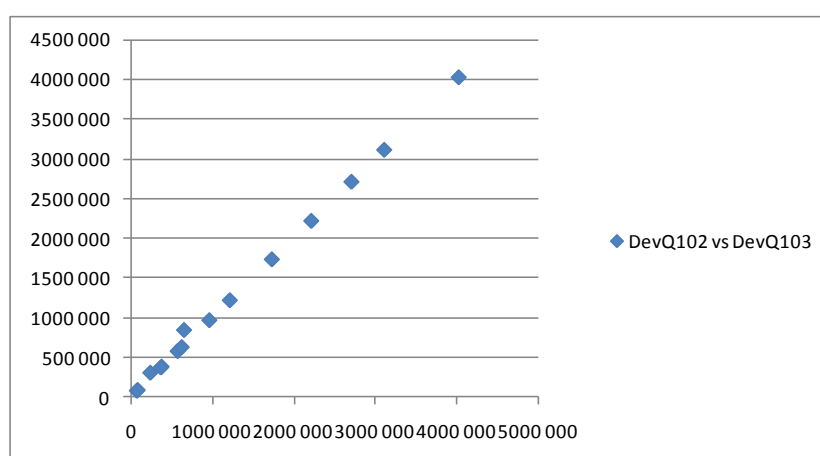


Figure 17*. Paid claims at the 102nd development quarter against the ones at the 103rd development quarter.

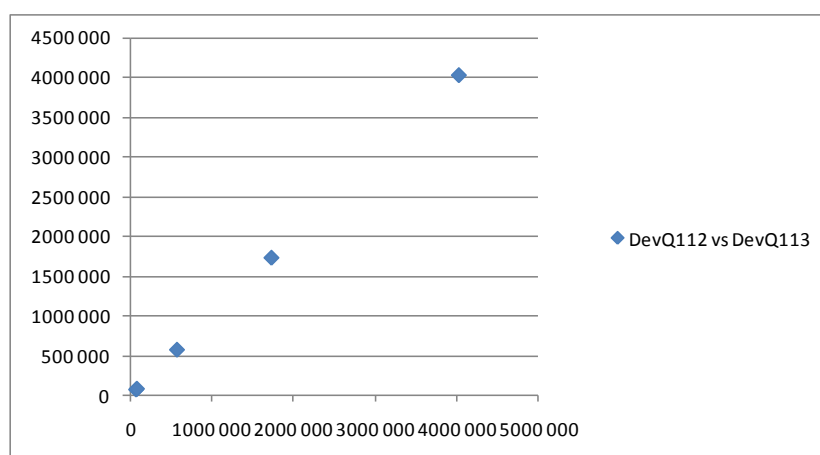


Figure 18*. Paid claims at the 112th development quarter against the ones at the 113th development quarter.

7.3.3 Appendix III.c – Incurred claims data variance assumption (CL3)

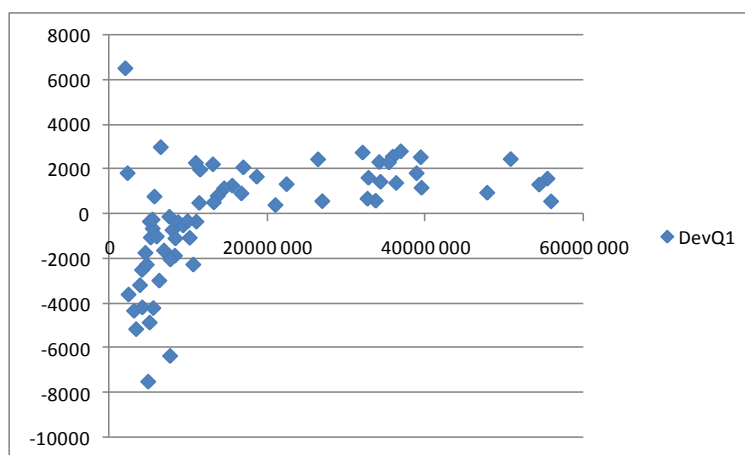


Figure 1**. Incurred claims at the 1st development quarter against the corresponding weighted residuals.

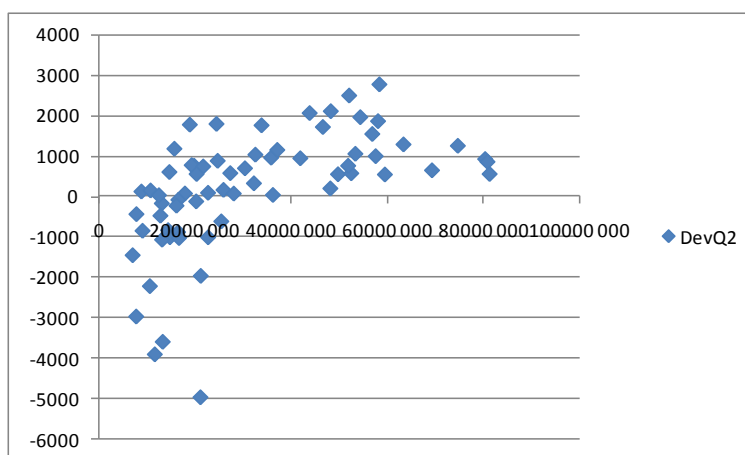


Figure 2**. Incurred claims at the 2nd development quarter against the corresponding weighted residuals.

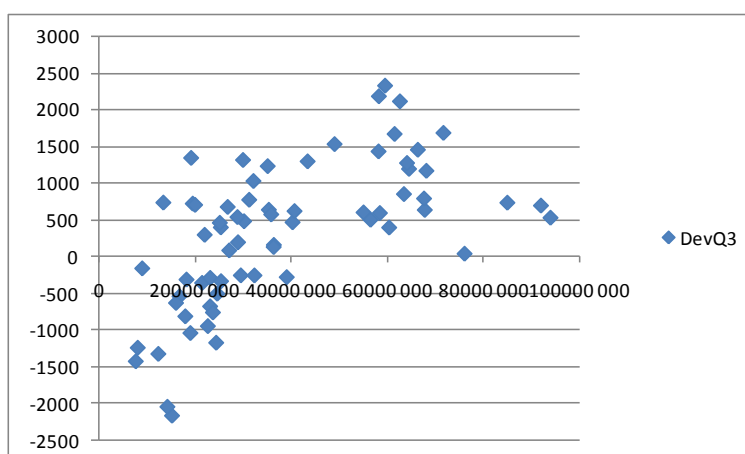


Figure 3**. Incurred claims at the 3rd development quarter against the corresponding weighted residuals.

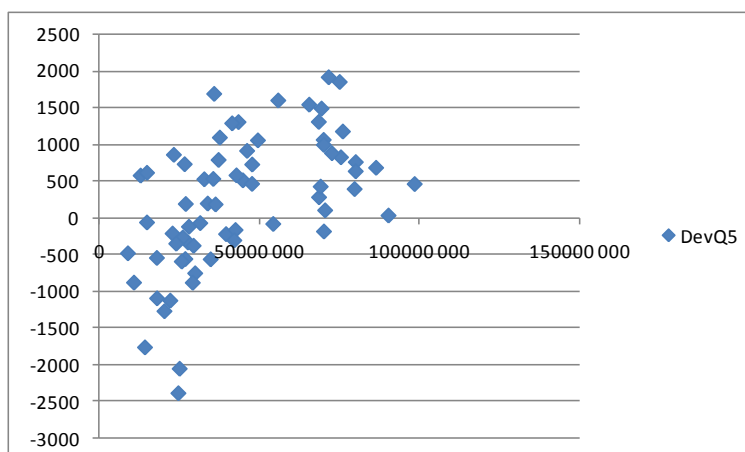


Figure 4**. Incurred claims at the 5th development quarter against the corresponding weighted residuals.

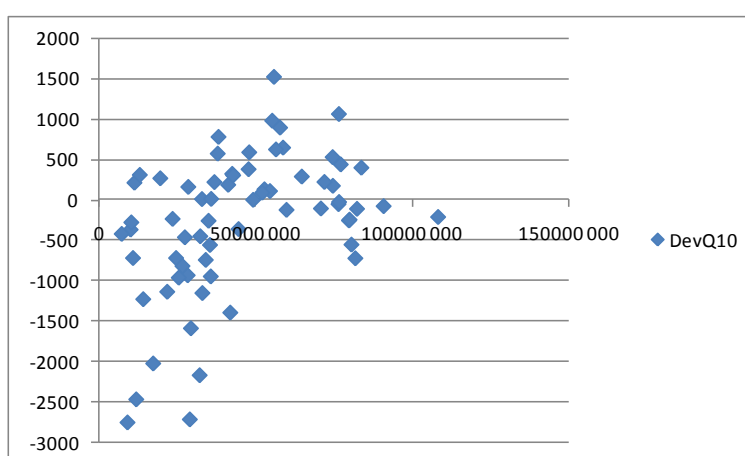


Figure 5**. Incurred claims at the 10th development quarter against the corresponding weighted residuals.

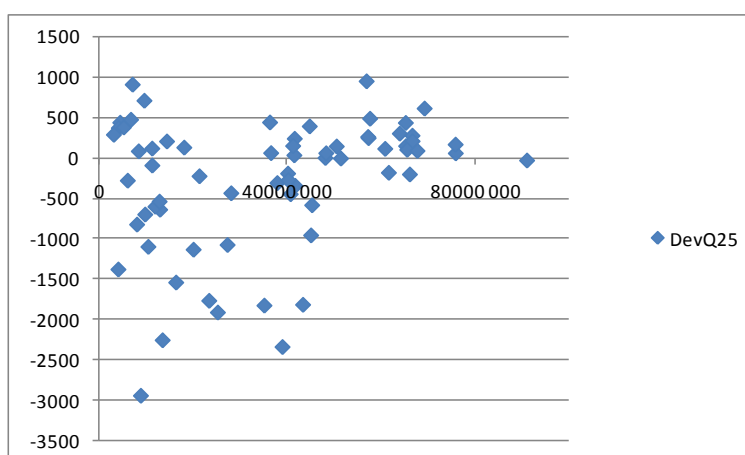


Figure 5**. Incurred claims at the 25th development quarter against the corresponding weighted residuals.

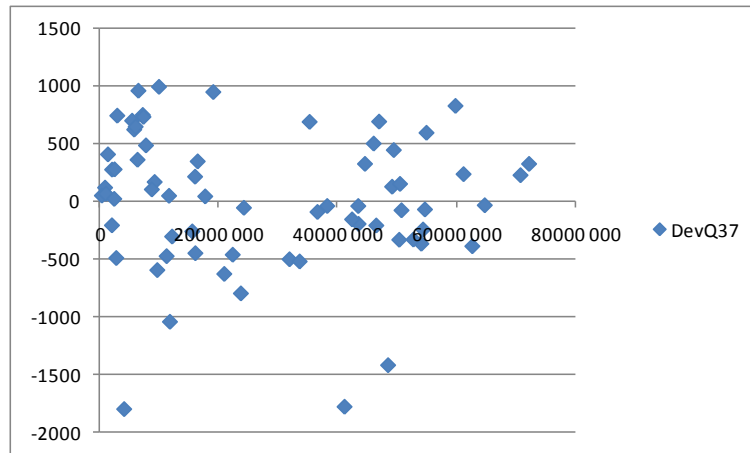


Figure 7**. Incurred claims at the 37th development quarter against the corresponding weighted residuals.

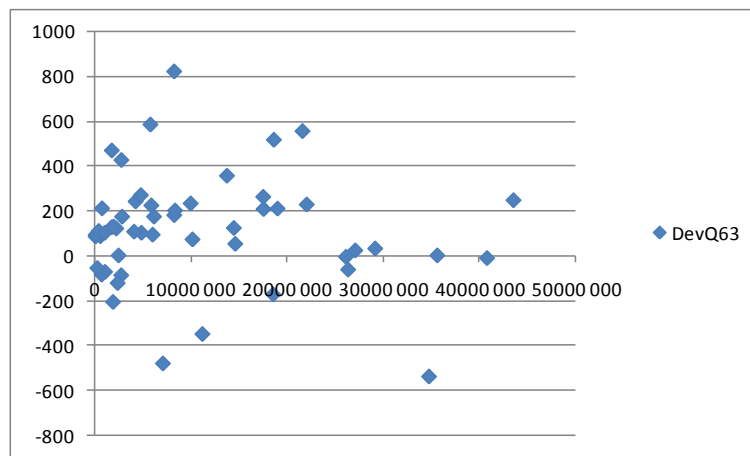


Figure 8**. Incurred claims at the 63rd development quarter against the corresponding weighted residuals.

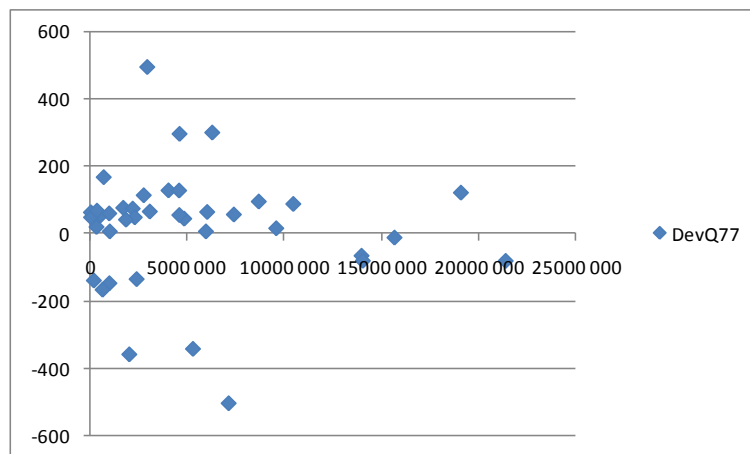


Figure 9**. Incurred claims at the 77th development quarter against the corresponding weighted residuals.

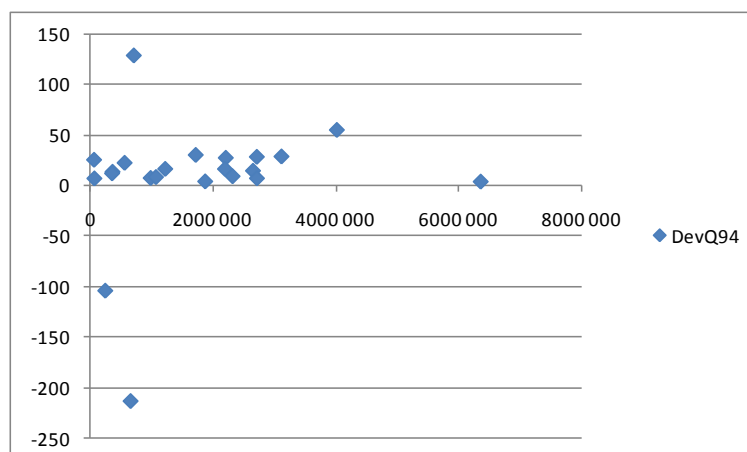


Figure 10**. Incurred claims at the 94th development quarter against the corresponding weighted residuals.

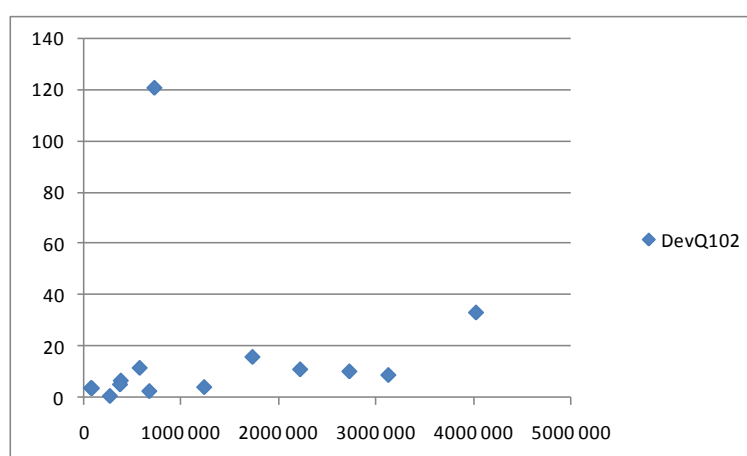


Figure 11**. Incurred claims at the 102nd development quarter against the corresponding weighted residuals.

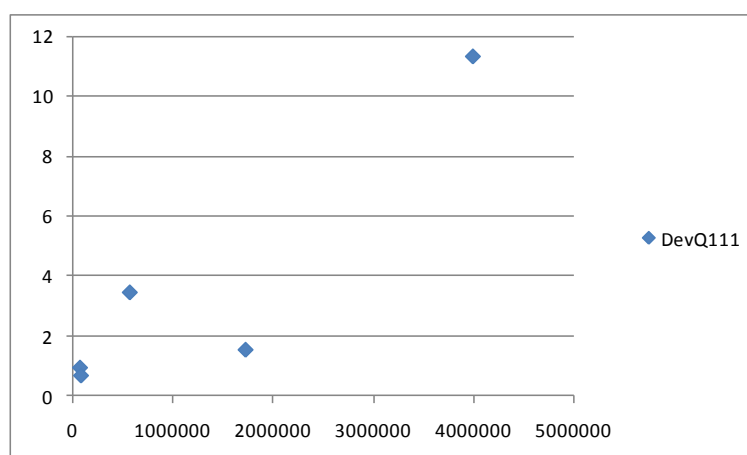


Figure 12**. Incurred claims at the 111th development quarter against the corresponding weighted residuals.

7.3.4 Appendix III.d – Paid claims data variance assumption (CL3)

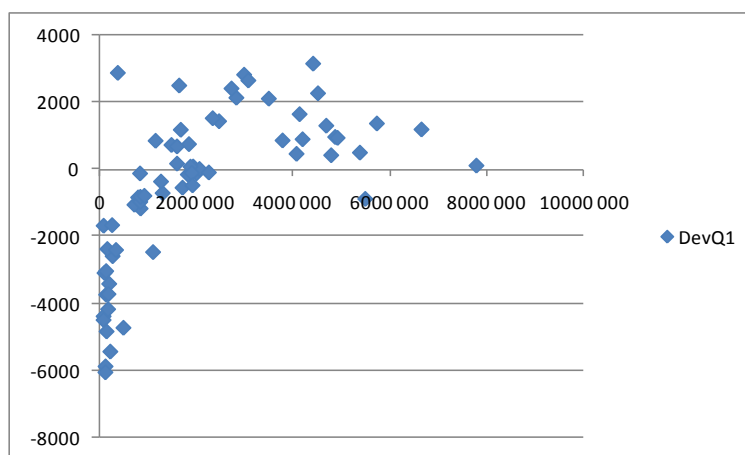


Figure 13**. Paid claims at the 1st development quarter against the corresponding weighted residuals.

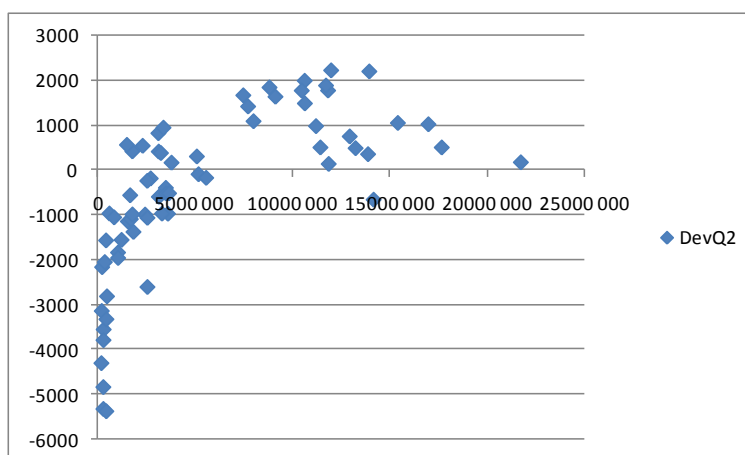


Figure 14**. Paid claims at the 2nd development quarter against the corresponding weighted residuals.

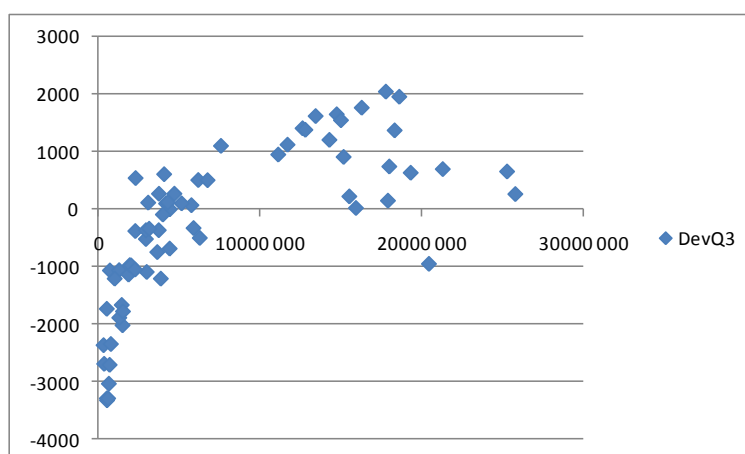


Figure 15**. Paid claims at the 3rd development quarter against the corresponding weighted residuals.

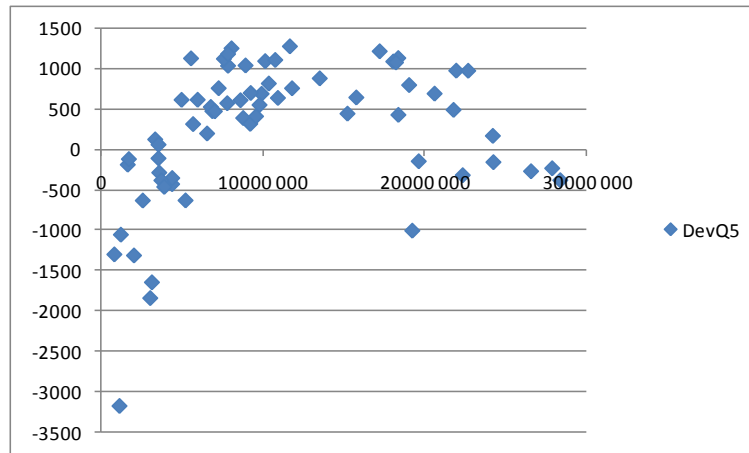


Figure 16**. Paid claims at the 5th development quarter against the corresponding weighted residuals.

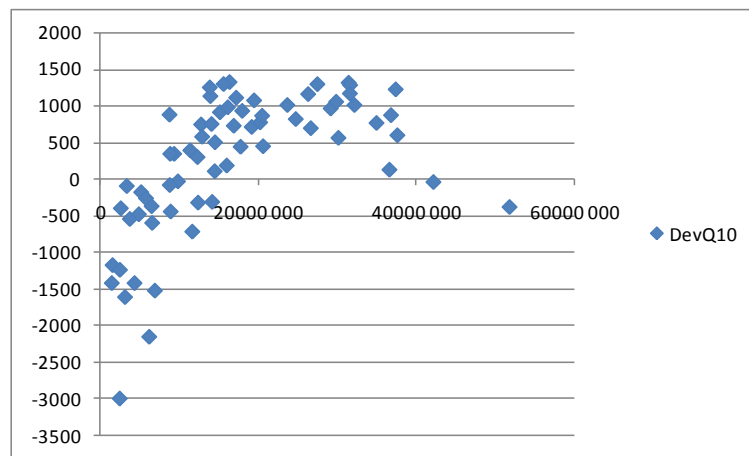


Figure 17**. Paid claims at the 10th development quarter against the corresponding weighted residuals.

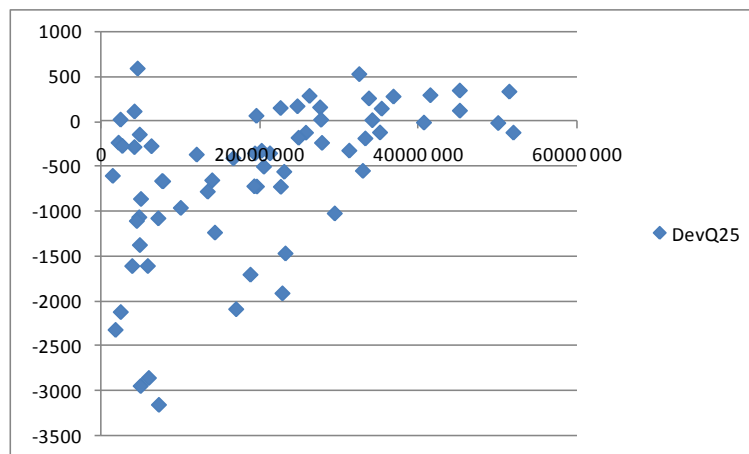


Figure 18**. Paid claims at the 25th development quarter against the corresponding weighted residuals.

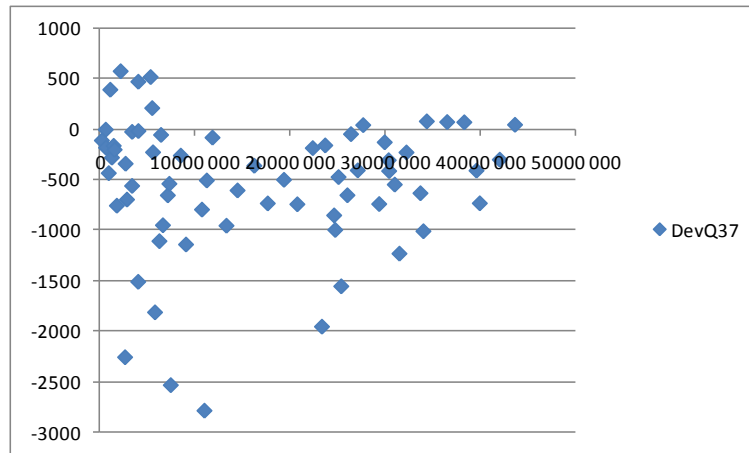


Figure 19**. Paid claims at the 37th development quarter against the corresponding weighted residuals.

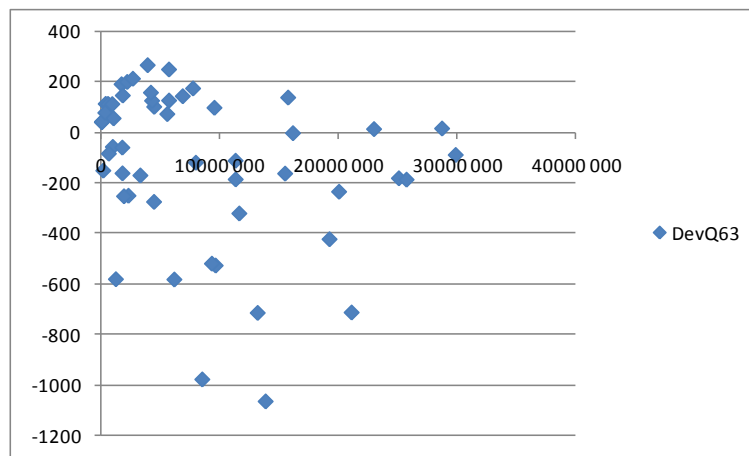


Figure 20**. Paid claims at the 63rd development quarter against the corresponding weighted residuals.

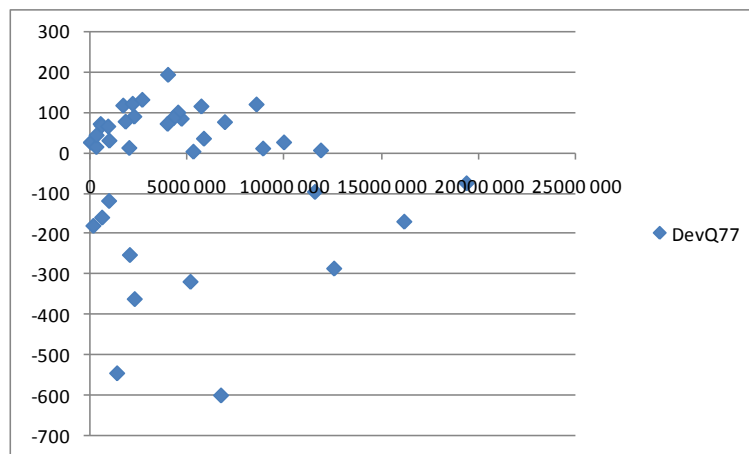


Figure 21**. Paid claims at the 77th development quarter against the corresponding weighted residuals.

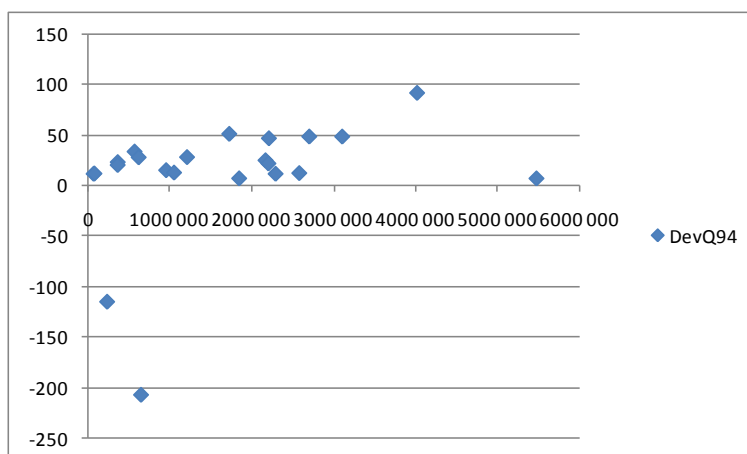


Figure 22**. Paid claims at the 94th development quarter against the corresponding weighted residuals.

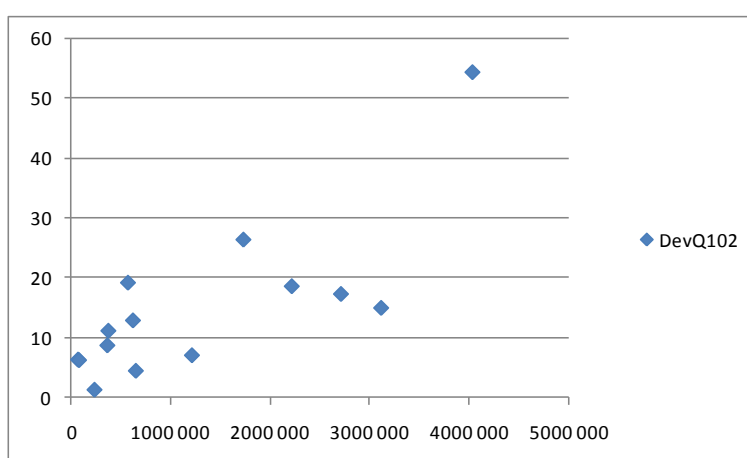


Figure 23**. Paid claims at the 102nd development quarter against the corresponding weighted residuals.

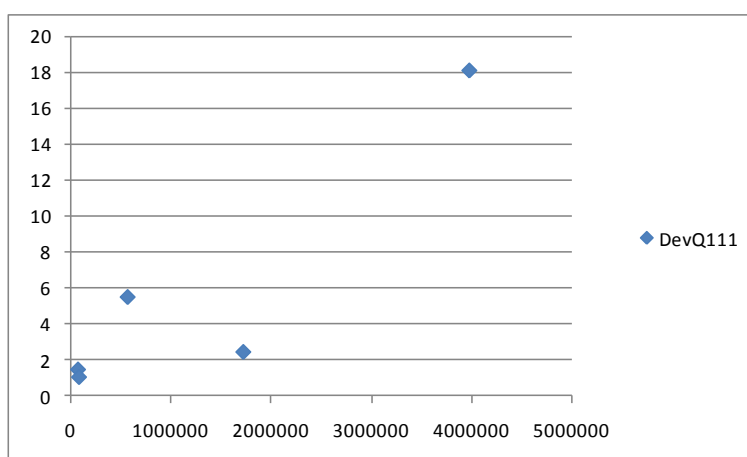


Figure 24**. Paid claims at the 111th development quarter against the corresponding weighted residuals.

7.4 Appendix IV. The Estimation Error of the BF method

In the figures below 1^{st} stands for the first bootstrapping approach, and 2^{nd} – for the second bootstrapping approach.

