

BOOTSTRAP METHODOLOGY IN CLAIM RESERVING

Paulo J. R. Pinheiro

João Manuel Andrade e Silva

Maria de Lourdes Centeno

ABSTRACT

In this article, we use the bootstrap technique to obtain prediction errors for different claim-reserving methods, namely, the chain ladder technique and methods based on generalized linear models. We discuss several forms of performing the bootstrap and illustrate the different solutions using the data set from Taylor and Ashe (1983), which has already been used by several authors.

INTRODUCTION

The prediction of an adequate amount to face the responsibilities assumed by an insurance company is a major subject in actuarial science. Despite its well-known limitations, the chain ladder technique (see for instance Taylor (2000) for a presentation of this technique) is the most widely applied claim-reserving method. Moreover, in recent years, considerable attention has been given to the discussion of possible relationships between the chain ladder and various stochastic models (Mack, 1993, 1994; Mack and Venter, 2000; Verrall, 1991, 2000; Renshaw and Verrall, 1994; England and Verrall, 1999; etc.).

The bootstrap technique has proved to be a very useful tool in many fields and can be particularly interesting to assess the variability of the claim-reserving predictions and to construct upper limits at an adequate confidence level. Some applications of the bootstrap technique to claim reserving can be found in Lowe (1994), England and Verrall (1999), and Taylor (2000).

The application of the bootstrap technique to claim reserving is not straightforward and, in our opinion, the applications found in the actuarial literature were not the most adequate.

Paulo J. R. Pinheiro is at Zurich Companhia de Seguros, SA. João Manuel Andrade e Silva and Maria de Lourdes Centeno are at CEMAPRE, ISEG, Technical University of Lisbon. This research was partially supported by Fundação para a Ciência e a Tecnologia (FCT) and by POCTI. We are grateful to two anonymous referees for their comments to an earlier version of this article.

FIGURE 1

Pattern of the Available Data

Origin	Development Year						
Year	1	2	...	j	...	$n-1$	n
1	C_{11}	C_{12}	...	C_{1j}	...	$C_{1,n-1}$	C_n
2	C_{21}	C_{22}	...	C_{2j}	...	$C_{2,n-1}$	
...	...	...	...	...			
i	C_{i1}	C_{i2}	...	$C_{i,n+1-i}$			
...							
$n-1$	$C_{n-1,1}$	$C_{n-1,2}$					
n	$C_{n,1}$						

The main issue to be treated in this article concerns the definition of the residuals to be considered in the bootstrap procedure. We also discuss two different methodologies of performing the bootstrap.

The problem of claim reserving can be summarized in the following way: given the available information about the past, how can we obtain an estimate of the future payments (or the number of claims to be reported) due to claims occurred in those years? Furthermore, we need to determine a prudential margin, which is to say, we want to estimate an upper limit for the reserve with an adequate level of confidence.

Let C_{ij} represent either the incremental claim amounts or the number of claims arising from accident year i and development year j and let us assume that we are in year n and that we know all the past information, i.e., C_{ij} ($i = 1, 2, \dots, n$ and $j = 1, 2, \dots, n+1-i$). The available data present a characteristic pattern, which can be seen in Figure 1. From now on, and without loss of generality, we consider that the C_{ij} are the incremental claim amounts.

More than to predict the individual values, C_{ij} ($i = 2, 3, \dots, n$ and $j = n+2-i, n+3-i, \dots, n$), we are interested in the prediction of the rows total, $C_{i\bullet}$ ($i = 2, 3, \dots, n$), i.e., the amounts needed to face the claims occurred in year i and especially in the aggregate prediction, C , which represents the expected total liability. Keep in mind that we want to obtain upper limits to the forecasts and to associate a confidence level to those limits.

In "Generalized Linear Models and Claim-Reserving Methods" we present a brief review of generalized linear models (GLM) and their application to claim reserving, whereas in "The Bootstrap Technique" we discuss the application of the bootstrap technique. In "An Application" we illustrate the two different methods presented in "The Bootstrap Technique" section to the data set provided in Taylor and Ashe (1983), which has already been used by several authors.

GENERALIZED LINEAR MODELS AND CLAIM-RESERVING METHODS

Following Renshaw and Verrall (1994) we can formulate most of the stochastic models for claim reserving by means of a particular family of generalized linear models (see

McCullagh and Nelder, 1989, for an introduction to GLM). The structure of those GLM will be given by

$$Y_{ij} \sim f(y; \mu_{ij}, \phi) \quad (1)$$

with independent Y_{ij} 's, $\mu_{ij} = E(Y_{ij})$, and where $f(\cdot)$, the density (probability) function of Y_{ij} , belongs to the exponential family. ϕ is a scale parameter;

$$\eta_{ij} = g(\mu_{ij}); \quad (2)$$

and

$$\eta_{ij} = c + \alpha_i + \beta_j, \quad (3)$$

with $\alpha_1 = \beta_1 = 0$ to avoid overparameterization.

Assumption (1) requires independent incremental claim amounts. This is a crucial assumption, which is often not fulfilled.

It is common in claim reserving to consider three possible distributions for the variable C_{ij} : lognormal, gamma, or Poisson. For models based on gamma or Poisson distributions, the relations (1)–(3) define a GLM with $Y_{ij} = C_{ij}$ denoting the incremental claim amounts. The link function is $\eta_{ij} = \ln(\mu_{ij})$.

When we consider that the claim amounts follow a lognormal distribution, see Kremer (1982), Verrall (1991), or Renshaw (1994), among others, we have that $Y_{ij} = \ln(C_{ij})$ has a normal distribution and consequently the relations (1)–(3) still continue to define a GLM for the logs of the incremental claim amounts. Now the link function is given by $\eta_{ij} = \mu_{ij}$ and the scale parameter is the variance of the normal distribution, i.e., $\phi = \sigma^2$.

The linear structure given by (3) implies that the estimates for some of the parameters depend on one observation only, i.e., there is a perfect fit for these observations. If the available data follow the pattern shown in Figure 1, it is straightforward to see that $\hat{\mu}_{1,n} = y_{1,n}$ and that $\hat{\mu}_{n,1} = y_{n,1}$.

When we define a GLM, we can omit the distribution of Y_{ij} and specify only the variance function and estimate the parameters by maximum quasi likelihood (McCullagh and Nelder, 1989) instead of maximum likelihood. The estimators remain consistent. In this formulation, we replace the distributional assumption by $\text{var}(Y_{ij}) = \phi V(\mu_{ij})$, where $V(\cdot)$ is called the variance function. As we know, for the normal distribution $V(\mu_{ij}) = 1$, for the Poisson distribution (or “overdispersed” when $\phi > 1$) $V(\mu_{ij}) = \mu_{ij}$ and for the gamma distribution $V(\mu_{ij}) = \mu_{ij}^2$.

It is well known that a GLM with the linear structure given by (3) and $V(\mu_{ij}) = \mu_{ij}$, i.e., an overdispersed Poisson distribution, gives the same predictions as those obtained by the chain ladder technique when we use a full triangle, as is the case in this article (see Renshaw and Verrall, 1994). However, if we use a quasi overdispersed Poisson, it is necessary to impose the constraint that the sum of incremental claims in each column is greater than zero. Note that a similar, but more complicated, constraint applies to quasi gamma models and that we need a stronger constraint for lognormal, gamma, or Poisson models.

As we said, in claim reserving, the figures of interest will be the aggregate value $Y_{\bullet} = \sum_{i=2}^n \sum_{j=n+2-i}^n Y_{ij}$ and the rows total $Y_{i\bullet} = \sum_{j=n+2-i}^n Y_{ij}$. The predicted values will be given by $\hat{\mu}_{\bullet} = \sum_{i=2}^n \sum_{j=n+2-i}^n \hat{\mu}_{ij}$ and $\hat{\mu}_{i\bullet} = \sum_{j=n+2-i}^n \hat{\mu}_{ij}$, respectively. To obtain these forecasts the procedure will be as follows:

- define the model,
- estimate the parameters c, α_i, β_j for $i, j = 1, 2, \dots, n$ and ϕ ,
- obtain the fitted values $\hat{\mu}_{ij}$ ($i = 1, 2, \dots, n$ and $j = 1, 2, \dots, n + 1 - i$),
- check the model,
- obtain the "individual" forecasts $\hat{\mu}_{ij} = \hat{c} + \hat{\alpha}_i + \hat{\beta}_j$ ($i = 2, \dots, n$ and $j = n + 2 - i, \dots, n$),
- obtain the forecasts for the rows reserve $\hat{\mu}_{i\bullet} = \sum_{j=n+2-i}^n \hat{\mu}_{ij}$ ($i = 2, \dots, n$), and
- obtain the forecast for the total reserve $\hat{\mu}_{\bullet} = \sum_{i=2}^n \hat{\mu}_{i\bullet}$.

Obtaining estimates for the standard error of prediction is a more difficult task. Renshaw (1994), using first degree Taylor expansions, deduced some approximations to the standard errors (see also England and Verrall, 1999). These values are, when the log link function is used, given by

- Standard error for the "individual" predictions:

$$\sqrt{E(Y_{ij} - \hat{\mu}_{ij})^2} \cong \sqrt{\text{var}(Y_{ij}) + \text{var}(\hat{\mu}_{ij})} \cong \sqrt{\phi V(\mu_{ij}) + \mu_{ij}^2 \text{var}(\hat{\eta}_{ij})}, \quad (4)$$

where $V(\cdot)$ is the variance function and $\text{var}(\hat{\eta}_{ij})$ is obtained as a function of the covariance matrix of the estimators and is usually available from most statistical software. The term μ_{ij}^2 is a consequence of the link function chosen.

- Standard error for the row totals:

$$\begin{aligned} \sqrt{E(Y_{i\bullet} - \hat{\mu}_{i\bullet})^2} &\cong \sqrt{\text{var}(Y_{i\bullet}) + \text{var}(\hat{\mu}_{i\bullet})} \\ &\cong \sqrt{\sum_j \phi V(\mu_{ij}) + \sum_j \mu_{ij}^2 \text{var}(\hat{\eta}_{ij}) + 2 \sum_{\substack{j_1, j_2 \\ j_2 > j_1}} \mu_{ij_1} \mu_{ij_2} \text{cov}(\hat{\eta}_{ij_1}, \hat{\eta}_{ij_2})}, \end{aligned} \quad (5)$$

where the summations are made for the "individual" forecasts in each row.

- Standard error for the grand total:

$$\sqrt{E(Y_{\bullet} - \hat{\mu}_{\bullet})^2} \cong \sqrt{\sum_{i,j} \phi V(\mu_{ij}) + \sum_{i,j} \mu_{ij}^2 \text{var}(\hat{\eta}_{ij}) + \sum_{\substack{i_1, j_1 \\ i_2, j_2 \\ i_1, j_1 \neq i_2, j_2}} \mu_{i_1 j_1} \mu_{i_2 j_2} \text{cov}(\hat{\eta}_{i_1 j_1}, \hat{\eta}_{i_2 j_2})}, \quad (6)$$

where the summations are made for all the "individual" forecasts.

Those estimates are difficult to calculate and are only approximate values, even in the hypothesis that the model is correctly specified. This is the main reason to take advantage of the bootstrap technique.

THE BOOTSTRAP TECHNIQUE

The bootstrap technique is a particular resampling method used to estimate, in a consistent way, the variability of a parameter. This resampling method replaces theoretical deductions in statistical analysis by repeatedly resampling the “original” data and making inferences from the resamples.

Presentation of the bootstrap technique could be easily found in the literature (see, for instance, Efron and Tibshirani, 1993; Shao and Tu, 1995; or Davison and Hinkley, 1997).

The bootstrap technique must be adapted to each situation. For the linear model (“classical” or generalized) it is common to adopt one of two possible ways:

- paired bootstrap—the resampling is done directly from the observations (values of y and the corresponding lines of the X matrix in the regression model); and
- residuals bootstrap—the resampling is applied to the residuals of the model.

Despite the fact that the paired bootstrap is more robust than the residual bootstrap, only the latter could be implemented in the context of the claim reserving, given the dependence between some observations and the parameter estimates.

To implement a bootstrap analysis we need to choose a model, to define an adequate residual and to use a bootstrap prediction procedure.

To define the most adequate residuals for the bootstrap, it is important to remember two points:

- the resampling is based on the hypothesis that the residuals are independent and identically distributed; and
- it is indifferent to resample the residuals or the residuals multiplied by a constant, as long as we take that fact into account in the generation of the pseudo data.

Within the framework of a GLM we could use different types of residuals (Pearson, deviance, Anscombe, etc.). In this article, our starting point will be the Pearson residuals defined by

$$r_{ij}^{(P)} = \frac{y_{ij} - \hat{\mu}_{ij}}{\sqrt{\widehat{\text{var}}(Y_{ij})}} = \frac{y_{ij} - \hat{\mu}_{ij}}{\sqrt{\hat{\phi} V(\hat{\mu}_{ij})}}. \quad (7)$$

Since ϕ is constant for the data set, we can take advantage of the second point and use

$$r_{ij}^{(P^*)} = \frac{y_{ij} - \hat{\mu}_{ij}}{\sqrt{V(\hat{\mu}_{ij})}} \quad (8)$$

instead of $r_{ij}^{(P)}$ in the bootstrap procedure, that is to ignore, at this stage, the scale parameter. When using a normal model it is trivial to see that these residuals are equivalent to the classical residuals, $y_{ij} - \hat{\mu}_{ij}$, since $V(\mu_{ij}) = 1$.

However, these residuals need to be corrected since the available data combined with the linear structure adopted in the model lead to some residuals of value 0 (as we have

already mentioned, in the typical case, $y_{1,n} = \hat{\mu}_{1,n}$ and $y_{n,1} = \hat{\mu}_{n,1}$). These residuals should not be considered as observations of the underlying random variable and consequently should not be considered in the bootstrap procedure.

As in the classical linear model (see Efron and Tibshirani, 1993), it is more adequate to work with the standardized Pearson residuals and not the Pearson residuals, since only the former could be considered as identically distributed. It is well known the standardized Pearson residuals are given by

$$r_{ij}^{(P^{**})} = \frac{r_{ij}^P}{\sqrt{1 - h_{ij}}}, \quad (9)$$

where the factor h_{ij} is the corresponding element of the diagonal of the “hat” matrix. For the “classical” linear model, this matrix is given by

$$H = X(X^T X)^{-1} X^T$$

and for a GLM it can be generalized using

$$H = X(X^T W X)^{-1} X^T W,$$

where W is a diagonal matrix with generic element given by

$$\left(V(\mu_{ij}) \left(\frac{\partial \eta_{ij}}{\partial \mu_{ij}} \right)^2 \right)^{-1}$$

(see McCullagh and Nelder, 1989).

Considering the structure of our models (log link and variance functions), this generic element will be given by $\mu_{ij}^{2-\kappa}$, with $\kappa = 1$ for the quasi overdispersed Poisson and $\kappa = 2$ for the quasi gamma model.

Note that similar procedures could be defined if we use another kind of residuals, namely, the deviance residuals.

Let us now briefly discuss the bootstrap prediction procedure. To obtain an upper confidence limit for the forecasts of the aggregate values we can use two approaches: the first one takes advantage of the Central Limit Theorem and consists on approximating the distribution of the reserve by means of a normal distribution with expected value given by the initial forecast (with the original data) and standard deviation given by the “standard error of prediction.” The main difference between the bootstrap estimation of these standard errors and the theoretical approximation stated in the “Generalized Linear Models and Claim-Reserving Methods” section is that we estimate the variance of the estimator by means of a bootstrap estimate instead of using the (approximate) theoretical expression. For a detailed presentation of this method (in a general environment) see Efron and Tibshirani (1993). England and Verrall (1999) use this approach in claim reserving and suggest a bias correction for the bootstrap

estimate to allow the comparison between the bootstrap standard error of prediction and the theoretical approximation presented in the “Generalized Linear Models and Claim-Reserving Methods” section. The reason for this correction is the fact that the variance of the residuals is smaller than the variance of the underlying random variable. Moreover, the variance of each residual depends not only on the random variable but also on the data structure of the model. The solution used by England and Verrall (1999) consists in the introduction of a global correction. However, when we use the residuals corrected by the h factor, we use a different correction for each residual to guarantee (assuming the framework of the model) that they have the same variance as the underlying random variable. So, the global correction should not be used when the residuals have already been corrected by the h factor (see Moulton and Zeger, 1991). The bootstrap standard error of prediction with bias correction will be given by

$$\begin{aligned} SEP_b(\mu) &= \sqrt{\hat{\phi} V(\hat{\mu}) + \frac{N}{N-p} (SE_b(\hat{\mu}))^2} \\ &= \sqrt{\hat{\phi} \sum \hat{\mu}_{ij}^\kappa + \frac{N}{N-p} (SE_b(\hat{\mu}))^2}, \end{aligned} \quad (10)$$

and without correction by

$$SEP_b(\mu) = \sqrt{\hat{\phi} \sum \hat{\mu}_{ij}^\kappa + (SE_b(\hat{\mu}))^2}, \quad (11)$$

where $\kappa = 1$ for the quasi overdispersed Poisson and $\kappa = 2$ for the quasi gamma model, μ stands for the row totals, μ_i ($i = 2, 3, \dots, n$), or the aggregate total, μ_i , and the summation is done over the adequate individual predicted values. $\hat{\phi}$ and $\hat{\mu}$ are quasi-maximum likelihood estimates of the corresponding parameters, N is the number of observations, p is the number of parameters (usually $N = n(n-1)$ and $p = 2n-1$), while $SE_b(\hat{\mu})$ is the bootstrap estimate of the standard error of the estimator $\hat{\mu}$, i.e.,

$$SE_b(\hat{\mu}) = \sqrt{\frac{1}{B} \sum_{k=1}^B (\mu_k^* - \hat{\mu})^2},$$

where B is the number of bootstrap replicates and μ_k^* is the bootstrap estimate of μ in the k th replicate ($k = 1, 2, \dots, B$). Note that, whereas $\hat{\mu}$ is obtained from the quasi-maximum likelihood equation (using (2) and (3)), $\hat{\phi}$ is a moment estimator based on the vector $y - \mu$, i.e.,

$$\hat{\phi} = \frac{1}{N-p} \sum_{i,j} \frac{(y_{ij} - \hat{\mu}_{ij})^2}{V(\hat{\mu}_{ij})}.$$

The second approach (see Davison and Hinkley, 1997) is more computer intensive, since it requires two resampling procedures in the same bootstrap “iteration,” but the results should be more robust against deviations from the hypothesis of the model. The idea is to define an adequate “prediction error” as a function of the bootstrap

FIGURE 2**First Bootstrap Procedure****Stage 1 – The preliminaries**

- Estimation of the model parameters c, α_i, β_j ($i, j = 1, 2, \dots, n$) and ϕ .
- Calculation of the fitted values, $\hat{\mu}_{ij}$ ($i = 1, 2, \dots, n$ and $j = 1, 2, \dots, n + 1 - i$).
- Calculation of the residuals $r_{ij} = \psi(y_{ij}, \hat{\mu}_{ij})$.
- Forecasts with the original data $\hat{\mu}_{ij}, \hat{\mu}_{i\bullet}$, and $\hat{\mu}_{\bullet}$ ($i = 2, \dots, n$ and $j = n + 2 - i, \dots, n$).

Stage 2 – Bootstrap loop (to be repeated B times)

- Resample the residuals obtained in stage 1 (original data) using replacement $\rightarrow r_{ij}^*$.
- Create the pseudo data y_{ij}^* , solving $r_{ij}^* = \psi(y_{ij}^*, \hat{\mu}_{ij})$.
- Estimate the model with the pseudo data and obtain the bootstrap forecast $\hat{\mu}_{ij}^*, \hat{\mu}_{i\bullet}^*$, and $\hat{\mu}_{\bullet}^*$.
- Keep the bootstrap forecasts $\hat{\mu}_{i\bullet}^{(b)} = \hat{\mu}_{i\bullet}^*$ and $\hat{\mu}_{\bullet}^{(b)} = \hat{\mu}_{\bullet}^*$, b being the index of the cycle.

Stage 3 – Bootstrap data analysis

- Obtain the bootstrap estimate for $\text{var}(\hat{\mu}_{i\bullet})$ and $\text{var}(\hat{\mu}_{\bullet})$ by means of the empirical variance of the corresponding B bootstrap estimates. If we use the uncorrected residuals, we must correct the bias of such estimates by multiplying them by a factor equal to $N/(N - p)$, N being the number of observations in the data triangle and p the number of parameters in the linear structure.
- Apply the theoretical expressions of the standard error of prediction and use those estimates.

TABLE 1**Available Data**

	1	2	3	4	5	6	7	8	9	10
1	357,848	766,940	610,542	482,940	527,326	574,398	146,342	139,950	227,229	67,948
2	352,118	884,021	933,894	1,183,289	445,745	320,996	527,804	266,172	425,046	
3	290,507	1,001,799	926,219	1,016,654	750,816	146,923	495,992	280,405		
4	310,608	1,108,250	776,189	1,562,400	272,482	352,053	206,286			
5	443,160	693,190	991,983	769,488	504,851	470,639				
6	396,132	937,085	847,498	805,037	705,960					
7	440,832	847,631	1,131,398	1,063,269						
8	359,480	1,061,648	1,443,370							
9	376,686	986,608								
10	344,014									

estimate and a bootstrap simulation of the future reality and to record the value of this prediction error for each bootstrap “iteration.” We use the desired percentile of this prediction error and combine it with the initial prediction to obtain the upper limit of the prediction interval.

Figure 2 presents the different stages of the first bootstrap procedure and Figure 3 presents the second bootstrap procedure.

AN APPLICATION

Let us consider the data from Taylor and Ashe (1983), which are presented in Table 1 in incremental form. As already said, this data set has been used by several authors and acts as a sort of benchmark for claims reserving methods. England and Verrall

FIGURE 3
Second Bootstrap Procedure

Stage 1 – The preliminaries

- Estimation of the model parameters c, α_i, β_j ($i, j = 1, 2, \dots, n$) and ϕ ;
- Calculation of the fitted values, $\hat{\mu}_{ij}$ ($i = 1, 2, \dots, n$ and $j = 1, 2, \dots, n + 1 - i$);
- Calculation of the residuals $r_{ij} = \psi(y_{ij}, \hat{\mu}_{ij})$;
- Forecasts with the original data $\hat{\mu}_{ij}, \hat{\mu}_{i\bullet},$ and $\hat{\mu}_{\bullet}$ ($i = 2, \dots, n$ and $j = n + 2 - i, \dots, n$).

Stage 2 – Bootstrap loop (to be repeated B times)

Sub stage 2.1 – Bootstrap estimates

- Resample the residuals obtained in stage 1 (original data) using replacement $\rightarrow r_{ij}^*$.
- Create the pseudo data y_{ij}^* , solving $r_{ij}^* = \psi(y_{ij}^*, \hat{\mu}_{ij})$.
- Estimate the model with the pseudo data and obtain the bootstrap forecast $\hat{\mu}_{ij}^*, \hat{\mu}_{i\bullet}^*,$ and $\hat{\mu}_{\bullet}^*$.
- Keep the bootstrap forecasts $\hat{\mu}_{i\bullet}^{(b)} = \hat{\mu}_{i\bullet}^*$ and $\hat{\mu}_{\bullet}^{(b)} = \hat{\mu}_{\bullet}^*$, b being the index of the cycle.

Sub stage 2.2 – Pseudo reality

- Resample again the residuals obtained in stage 1 and select (with replacement) as many values as there are “individual” forecasts to be done $\rightarrow r_{ij}^{**}, (i = 2, \dots, n$ and $j = n + 2 - i, \dots, n)$.
- Create the pseudo reality, y_{ij}^{**} , solving $r_{ij}^{**} = \psi(y_{ij}^{**}, \hat{\mu}_{ij}) (i = 2, \dots, n$ and $j = n + 2 - i, \dots, n)$. $\hat{\mu}_{ij}$ are the predictions obtained in stage 1.
- Obtain the prediction errors $r_{i\bullet}^{(b)} = \psi(y_{i\bullet}^{**}, \hat{\mu}_{i\bullet}^*)$ and $r_{\bullet}^{(b)} = \psi(y_{\bullet}^{**}, \hat{\mu}_{\bullet}^*)$ and keep them.
- Return to the beginning of stage 2 until the B repetitions are completed.

Stage 3 – Bootstrap data analysis

- Use the percentile $k\%$ of the bootstrap observations of prediction error, for instance $r_{\bullet,k}^*$ for the grand total, and obtain the corresponding percentile of the provisions by solving $r_{\bullet,k}^* = \psi(y_{\bullet,k}^*, \hat{\mu}_{\bullet})$. $\hat{\mu}_{\bullet}$ is the prediction with the original data (stage 1).

(1999) have summarized the main results obtained by those authors and they also use bootstrap technique to evaluate the predictions standard errors related to the chain ladder approach. For that purpose, they consider a model with a quasi overdispersed Poisson data to define the residuals, taking advantage of the fact that this particular GLM generates in this particular situation the same estimates as the chain ladder technique. They use Pearson residuals without correction (given in relation (8)) and they follow the first procedure for the bootstrap, based on the estimation of a standard error of prediction. Despite that they use a GLM for the residual definition they obtain the predictions by means of the chain ladder. This difference is relevant since, in that particular example, the two methods do not agree for some pseudo data sets generated by the bootstrap. In fact, for some pseudo data sets we could have negative values in the northeast corner, which do not allow the use of this GLM. We will follow the same approach as England and Verrall and will call it mixed model, since the estimates are obtained by the chain ladder but the residuals are based on a quasi Poisson model.

Our purpose is to use this data set and analyze three issues: first, we correct the bootstrap procedure used by England and Verrall, using the h corrected residuals. Second, we illustrate the use of the alternative bootstrap procedure. Since the residuals, even corrected, could inherit the skewness of the original data, the usual bootstrap procedure could be misleading as we use an approximation of the normal distribution.

TABLE 2
Quasi Overdispersed Poisson Model

Year	Estimated Reserve	Without Correction		Corrected Residuals	
		SEP	Upper 95%	SEP	Upper 95%
2	94,634	108,949	273,840	110,936	277,108
3	469,511	216,284	825,266	213,571	820,804
4	709,638	258,377	1,134,631	257,996	1,134,003
5	984,889	304,002	1,484,928	301,476	1,480,772
6	1,419,459	376,754	2,039,163	370,270	2,028,499
7	2,177,641	488,362	2,980,925	498,900	2,998,258
8	3,920,301	792,406	5,223,693	771,798	5,189,795
9	4,278,972	1,081,289	6,057,533	1,029,730	5,972,726
10	4,625,811	2,034,469	7,972,214	2,039,736	7,980,877
Total	18,680,856	2,993,352	23,604,480	2,915,885	23,477,058

In such situations, it seems preferable to take advantage of the alternative bootstrap procedure. We will illustrate this procedure and analyze the differences in the upper limits for a confidence level of 95 percent with this data set. Third, we analyze the consequences of the choice of an alternative model in the framework of GLM. To illustrate this point, we consider the bootstrap predictions obtained by the gamma model and analyze how far they are from those obtained with the Poisson model. Notice that for this data set, the gamma model presents a clear advantage: all the pseudo data generated by this model allow their estimation by the same model. In all the bootstrap applications we have used $B = 1000$.

To discuss the first point we consider the same methodology as England and Verall, which will be called mixed model (quasi overdispersed Poisson for the residual definition and chain ladder for the predictions) and the first bootstrap procedure. We use the Pearson residuals without corrections and the residuals with the zeros corrected and the factor h (see Equation (9)).

Table 2 presents the standard errors of prediction for the two situations considered as well as the upper limits for a confidence level of 95 percent. Remember that the estimated reserve is the same.

As we can see, the standard errors of prediction obtained with the corrected data are not very different from those obtained with the uncorrected residuals (between -2 and 5 percent when we compare with the corrected ones). When we look to the upper limits the differences have the same way but the figures are necessarily smaller (between -1.2 and 1.4 percent). This result is acceptable if we remember that, when we use the uncorrected residuals, we apply a global correction. Note, however, that the use of the corrected residuals is more in accordance with the bootstrap theory (see Davison and Hinkley, 1997; Efron and Tibshirani, 1993; or Moulton and Zeger, 1991).

The second point is to compare the two bootstrap approaches. As we said before, the residuals, even corrected, could inherit the skewness of the data and, consequently, the second bootstrap procedure seems preferable, namely, where we face a significant

FIGURE 4
Bootstrap Forecast for Total and Year 3 Predictions

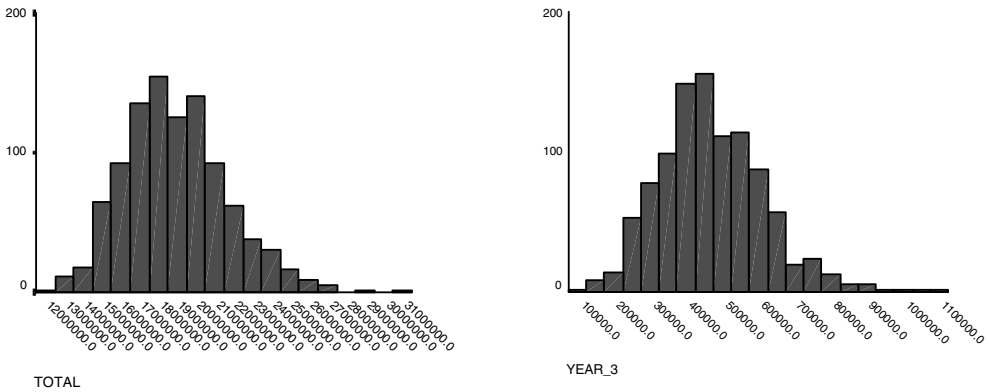


TABLE 3
SEP Against PPE Bootstrap Approaches—Overdispersed Poisson Model

Year	Estimated Reserve	Corrected Residuals	
		SEP—95%	PPE—95%
2	94,634	277,108	
3	469,511	820,804	886,168
4	709,638	1,134,003	1,175,163
5	984,889	1,480,772	1,520,295
6	1,419,459	2,028,499	2,106,503
7	2,177,641	2,998,258	3,085,471
8	3,920,301	5,189,795	5,286,592
9	4,278,972	5,972,726	6,215,378
10	4,625,811	7,980,877	9,370,058
Total	18,680,856	23,477,058	23,678,710

skewness. In this example, the skewness is not too high (the skewness coefficient of the corrected residuals is 0.437 with a standard error of 0.327).

To illustrate the effects of skewness, Figure 4 shows a histogram of the bootstrap forecasts (stage 2 in Figure 2 or stage 2.1 in Figure 3) for the aggregate prediction and for year 3 prediction. As it is expected, the skewness is much heavier in the results for year 3 than for the aggregate values.

To analyze the differences between the two approaches, we consider again the mixed model and obtain the upper limits using the two bootstrap procedures: SEP, that is the approach based on the standard error of prediction, against PPE, that is the other procedure, to obtain the adequate percentile of the prediction error. Table 3 presents the results for the upper 95 percent limit. Since some generated pseudo data set present negative values in the north-east corner, it is not possible to obtain the upper limit for year 2 when using PPE bootstrap procedure.

TABLE 4
Overdispersed Poisson Against Gamma Model—Theoretical Approximation

Year	Overdispersed Poisson			Gamma Model		
	Estimated Reserve	SEP	Upper 95%	Estimated Reserve	SEP	Upper 95%
2	94,634	110,258	275,992	93,316	46,505	169,810
3	469,511	216,265	825,235	446,504	165,315	718,423
4	709,638	261,114	1,139,132	611,145	182,889	911,971
5	984,889	303,822	1,484,632	992,023	262,013	1,422,996
6	1,419,459	375,374	2,036,894	1,453,085	361,748	2,048,107
7	2,177,641	495,911	2,993,342	2,186,161	541,888	3,077,486
8	3,920,301	791,169	5,221,658	3,665,066	969,223	5,259,294
9	4,278,972	1,048,624	6,003,804	4,122,398	1,210,801	6,113,988
10	4,625,811	1,984,733	7,890,405	4,516,073	1,716,813	7,339,978
Total	18,680,856	2,951,829	23,536,181	18,085,772	2,782,816	22,663,092

The second bootstrap procedure generates higher values for all the occurrence years. The differences amongst the results obtained with the two bootstrap procedures are more important for the first and the last years (namely, year 10). Note that the difference for the overall reserve is less than 1 percent but that this difference is higher for each individual year.

One advantage of the GLM is that we can extend this approach to other variance functions, that is, for instance, we can assume that the variance is proportional to the square of the mean instead of being proportional to the mean. Let us now compare the results for the quasi overdispersed Poisson against the gamma model. Table 4 shows the estimated reserve as well as the theoretical approximation to the standard error of prediction given by relations (5) and (6). Combining these two estimates and using the normal distribution we obtain the upper 95 percent confidence limits, which are also presented.

Two main conclusions can be drawn for this data set as follows.

- In our example, the gamma model produced smaller estimated reserves but the figures are not very different. The bigger difference is observed for year 4, where the value obtained with the gamma model is 14 percent less than those estimated with the overdispersed Poisson model. For the global prediction the same ratio is –3 percent.
- However the standard errors of prediction are quite different and consequently the estimated upper limits. These differences tend to be greater in the first years (estimation based on few predictions). The upper confidence limits present smoother differences, since they combine the standard errors of prediction with the estimated reserves. The upper limit for the global prediction is 4 percent smaller with the gamma model than with the overdispersed Poisson model.

Finally we compare the results obtained with the different models and the two bootstrap procedures. Table 5 presents those results when the residuals are corrected.

TABLE 5
Bootstrap Results (Corrected Residuals)

Year	Poisson Based			Gamma		
	Estimated Reserve	SEP—95%	PPE—95%	Estimated Reserve	SEP—95%	PPE—95%
2	94,634	277,108		93,316	168,108	224,222
3	469,511	820,804	886,168	446,504	712,166	797,805
4	709,638	1,134,003	1,175,163	611,145	906,906	996,543
5	984,889	1,480,772	1,520,295	992,023	1,430,559	1,522,673
6	1,419,459	2,028,499	2,106,503	1,453,085	2,041,856	2,117,230
7	2,177,641	2,998,258	3,085,471	2,186,161	3,066,776	3,240,837
8	3,920,301	5,189,795	5,286,592	3,665,066	5,285,036	5,649,816
9	4,278,972	5,972,726	6,215,378	4,122,398	6,134,969	7,063,204
10	4,625,811	7,980,877	9,370,058	4,516,073	7,364,444	9,911,301
Total	18,680,856	23,477,058	23,678,710	18,085,772	22,722,775	23,460,724

The main conclusion is that the aggregate prediction is much more influenced by the chosen model when the SEP bootstrap procedure is used (the Poisson upper limit is 3.3 percent higher than the gamma limit for the chosen confidence level) than when the PPE procedure is used (the difference is now 0.9 percent). This point is in accordance with the idea that the PPE bootstrap procedure is more robust. Nevertheless the differences are much more significant when we look at the results for each year and it is not clear that the PPE procedure produces upper limits more similar than the SEP. The choice of a particular model remains the main issue.

REFERENCES

- Davison, A. C., and D. V. Hinkley, 1997, *Bootstrap Methods and their Application*, Cambridge Series in Statistical and Probabilistic Mathematics (Cambridge: Cambridge University Press).
- Efron, B., and R. J. Tibshirani, 1993, *An Introduction to the Bootstrap* (London: Chapman and Hall).
- England, P., and R. Verrall, 1999, Analytic and Bootstrap Estimates of Prediction Errors in Claim Reserving, *Insurance: Mathematics and Economics*, 25: 281-293.
- Kremer, E., 1982, IBNR Claims and the Two Way Model of ANOVA, *Scandinavian Actuarial Journal*, 47-55.
- Lowe, J., 1994, A Practical Guide to Measuring Reserve Variability Using: Bootstrapping, Operational Time and a Distribution Free Approach, Presented at the 1994 General Insurance Convention, Institute of Actuaries and Faculty of Actuaries.
- Mack, T., and G. Venter, 2000, A Comparison of Stochastic Models that Reproduce Chain Ladder Reserve Estimates, *Insurance: Mathematics and Economics*, 26: 101-107.
- Mack, T., 1994, Which Stochastic Model is Underlying the Chain Ladder Model? *Insurance: Mathematics and Economics*, 15: 133-138.
- Mack, T., 1993, Distribution Free Calculation of the Standard Error of Chain Ladder Reserve Estimates, *ASTIN Bulletin*, 23(2): 213-225.

- McCullagh, P., and J. A. Nelder, 1989, *Generalized Linear Models*, 2nd edition (London: Chapman and Hall).
- Moulton, L. H., and S. L. Zeger, 1991, Bootstrapping Generalized Linear Models, *Computational Statistics and Data Analysis*, 11: 53-63.
- Renshaw, A. E., 1994, On the Second Moment Properties and the Implementation of Certain GLIM Based Stochastic Claims Reserving Models, Actuarial Research Papers no. 65, Department of Actuarial Science and Statistics, City University, London.
- Renshaw, A. E., and P. Verrall, 1994, A Stochastic Model Underlying the Chain Ladder Technique, Presented at the XXV ASTIN Colloquium, Cannes.
- Shao, J., and D. Tu, 1995, *The Jackknife and Bootstrap*, Springer Series in Statistics (Berlin: Springer-Verlag).
- Taylor, G., 2000, *Loss Reserving—An Actuarial Perspective* (Boston: Kluwer Academic Press).
- Taylor, G., and F. R. Ashe, 1983, Second Moments of Estimates of Outstanding Claims, *Journal of Econometrics*, 23: 37-61.
- Verrall, R., 2000, An Investigation into Stochastic Claims Reserving Models and the Chain-Ladder Technique, *Insurance: Mathematics and Economics*, 26: 91-99.
- Verrall, R., 1991, On the Estimation of Reserves from Loglinear Models, *Insurance: Mathematics and Economics*, 10: 75-80.

Copyright of Journal of Risk & Insurance is the property of Blackwell Publishing Limited and its content may not be copied or emailed to multiple sites or posted to a listserv without the copyright holder's express written permission. However, users may print, download, or email articles for individual use.