

DETERMINATION OF SAMPLE SIZE FOR AUDITING MEDICAL INSURANCE CLAIMS

Paritosh Chandra Sinha*, Santanu Kumar Ghosh**

Abstract *In the case of sample audit of the medical insurance claims, a complex as well as serious problem arises with the reality of the population characteristics of population error. The present study examines whether sample claims (those the auditors consider by nature or classification) is random or not. The paper puts forward that the sample size varies with respect to identification of different stages of audit by the Trainee, Processor, and Sr. Processor. For the first place that is, a sample audit by the Trainee, a firm's dependence on sample audit procedure should be taken into an utmost care. The study determines sample size econometrically through the application of the concept in Poisson distribution. Apart from these, the study suggests some ad-hoc population ranges and sample size against these ranges.*

Keywords: *Sample Audit, Sample Size, Population Error, and Medical Insurance Claims.*

With regard to financial appraisal of any project in the case of auditing medical insurance claims, most of the auditors prefer a *sample-audit* to a *total-audit* of the claims if the population size tends to become larger enough. The former audit procedure aims at cost reduction at the risk of non-audit of a larger part of the population generally. The desired size for a sample audit is obtained either following the Bayesian Method or the Precision Method of sample size determination (Kothari, 2004; p. 174-181). The former method however has a many advantages in the context of sample data analysis over the traditional hypothesis testing method (Kruschke, 2010; and Adcock, 1997; 1992). The procedure recommends that size of sample (n for every possible values of n) out of a population where the expected net gain, which is the difference of expected value of sample information over the cost of taking a sample, has become maximum. The method is tedious and cumbersome due to numerous subjective features. It lacks customization since the size limit could not be anticipated on *a-priori* basis for all of the unknown required information (see also, Gelman, 2008).

In contrast, the precision method is free from the earlier shortcomings and is subject to a few statistical assumptions. The medical claims in interest are assumed to appear randomly to an independent auditor. The said method determines the sample-size at a certain level of confidence for a given level of precision. Following the Precision Method in Kothari (2004), we determine the expected sample size in auditing the medical insurance claims of our client firm, ABC (name is changed to avoid leakages of client's internal information) Pvt. Solution LTD.

In the audit-process, the client firm firstly retrieves data of outsourced medical insurance claims electronically from its parent company in the USA, then encodes them for each of the claimants, and finally considers a random sampling technique for the purpose of sample audit in the query. The parent company in the USA receives the claims either directly from the corporate which provide medical insurance benefits for their employees and / or from the individuals who file their health care insurance policy by themselves. The client firm uses to audit the medical claims data by three groups of internal auditors viz. Trainee, Processor, and Senior Processor (hereinafter, Sr. Processor). The medical insurance claims which are assigned for the auditors depend on their audit skills, and a threshold limit in the dollar value of the claims. All claims up to the threshold limit of \$100 are audited by the Trainee only while those within the threshold limits of \$100 and \$300 interval are audited by the Processor. The medical claims of threshold limit above \$300 are audited by the Sr. Auditors. The auditors require auditing information in the different lines and the amount as well in the medical claims copy. The client firm in interest now requires to draw inference on appropriate sample size for three audit references viz. the lines in an insurance claim (that is, *Line Audit*), the amount in an insurance claim (that is, *Amount Audit*), and the number of medical claims (that is, *Claim Audit*). Now, besides determining the sample size in this paper, we investigate how far their sample sizes are precise.

We firstly test the nature of claims whether random or not, and then find out a feasible sample size. Since an audit seeks to reveal a true and fair view of the claims, a sample

* Research Scholar, The University of Burdwan, West Bengal, India; and Assistant Professor in Commerce, Rabindra Mahavidyalaya, Hooghly, W.B., India

** Professor, Department of Commerce, The University of Burdwan, West Bengal, India.

audit of the claims instead of a total audit entails risks in auditing process. These risks are of two folds - compliance due to non-audit of the claims and that due to acceptance of default claims (see also, Pathak, et al., 2005). In order to fulfill the first objective, we now examine a large sample data of 8733 claims over 27 classes of audit types and then, we use a comparatively medium sample data of 2125 claims in determining the sample-size. For the second purpose, we determine the sample size for three auditors' viz. Trainee, Processor and Sr. Processor with 95% level of confidence at 98% level of precision that is, at 2% level of error. The related issues are as follows.

THE NATURE OF SAMPLE DISTRIBUTION

Since the claims are selected for auditing on a random basis, inclusion of a claim into the sample follows a binomial distribution. The odds and the evens for collecting a claim as a sample variable are equi-probable. Now, since the respective claimants of the insurance claims seek to claim their funds with lesser hazard, then the number of default claims would be at minimum. The basic assumption is that the claimants are conservative on utilization of their hard earned money. Thus, even though the inclusion of a random claim in the sample follows binomial distribution, the concerned variables those represent default claims would follow *Poisson* distribution, where the probability of identifying a default claim is very low, but not zero. However, we would explain and explore this nature of the sample distribution accordingly in the section of Methodology.

Proper Sample Size

The number of default claims in any sample is random. The sample proportion of the default claims thus varies randomly with the variation in sample size. Now, with increase in the sample size, the sample proportion may either decrease or increase. In both the cases, the sampling technique may be biased with respect to the two extremes: (i) the samples are purely default, and (ii) the samples are purely non-default. Thus, with increase in the sample size, the sample default proportion would approach towards the population default proportion. The default proportion in representative sample drawn from the population therefore would not significantly differ from the population default proportion.

Level of Confidence

Now, in identifying a sample the auditor would always seek to provide for enough information about the population parameters. An auditor of our client firm ABC Solutions Pvt. Ltd. has to take care about the compliance due to non-audit of claims and /or compliance due to mistakes in audit process. This sort of errors persist in the audit process is termed

as Type: I error (α). If a randomly selected sample error differs from its true population error significantly, then we can extend our estimation to the extent of a true population error by allowing us towards a larger sample size. Thus, the level of confidence will indicate how far a predetermined precession level is covered with a certain sample size given that exact population proportion is true and known as well. Thus the default proportion with a larger sample size may provide approximate default proportion of the population. The magnitude of α contributes to the desired level of confidence ($1 - \alpha$).

Level of Precession

Since the data generation (i.e., inclusion of data as insurance claim) follows a systematic process in our client audit firm, the compliance due to acceptance of default claims is purely caused by non-audit of total population. The pre-determined value of the population error information may directly be generated from wrongly selected sample, and thereby, has provided inaccurate estimation of default claims about the population. These errors statistically warn rejection of true random sample and this is known as type-two error (β) statistically. The extreme value of the sample error for a randomly selected larger sample tends to be equal to the population error. Thus, a hypothetical 'total audit' for a known population size may result 100% precession, that is, zero percent type-two error. If the sample size is large and unbiased (that is, it provides a fair estimation of the population proportion), then by way of a sample audit we can estimate a reasonable level of precession ($1 - \beta$). Since the random variable, default proportion in our project follows *Poisson* distribution we require to examine the project of determination of sample size following the methodology described in the bellow.

METHODOLOGY

We assume that the parent population distribution of the medical insurance claim is normal. A claim at each trail in the auditing process of the firm either includes it or excludes it into the sample and the random claims follow binomial distribution. The sample claims would follow a normal distribution if its number of trails is infinitively large. Hence, the probability of identifying a default claim, which is behaviorally assumed to follow the *Poisson* distribution, would follow normal distribution if the sample size is large. In the followings, we explain the presumption about the *Poisson* framework. We then justify whether the claims those are included in each of the samples are random in nature or not.

Now our basic assumption about the insurance claims is that the claims are unbiased and drawn randomly from the population. We adopt the following methodology to establish

whether the claims are random in nature or not. In this regard, we utilize the reference of *claim Amount* with the both total value (in \$ amount) audited and error value (in \$ amount) against each of the claims over the different claim-types for a large sample of 8733 claims. The testable hypothesis is that the default proportion of each type (i.e., the proportion of error value to claim amount) is not significantly different from zero. For the different cases of insurance claim-types, the testable hypothesis is that the default proportions are not significantly different among the claims. The different claim types however are *vision, travel expenses, OTC, medical, licensed day care, dental orthodontia, dental order, and others* include 18 types of medical claims. An aggregate measure of the default claim, viz., f_l is given in proportion in the identity (1) as follows where a conservative error proportion under a claim type is given.

$$\partial_l = \frac{\sum_{k=1}^N ev_k}{\sum_{k=1}^N tv_k} \quad (1)$$

Here, in any sample of N size, ev_k and tv_k are *error value* and *total claim value* in a k^{th} claim ($k = 1, 2, 3 \dots N$) under an l -th type of claim. The variance statistic for l -th type claims data, $Var(f_l)$ for the default proportion f_l is defined in the identity (2) as follows where the default proportion in the k^{th} claim under l^{th} type of insurance is defined as:

$$\left(\frac{ev_k}{tv_k}\right)$$

$$Var(f_l) = Var(f_{lk}) = Var(f_k)_l = Var\left[\left(\frac{ev_k}{tv_k}\right)\right]_l \dots \dots \dots (2)$$

The findings in Table 1 suggest that the default proportion of the major claim types ranges within between the limits of 0.00030054 and 0.00979912 viz. for vision, travel expenses, OTC, medical, licensed day care, dental orthodontia, and dental order. The overall default proportion for these claim types excluding the other 18 claim types is 0.0041701, which is an efficient estimate of the default proportion with reference to *amount audit* in any sample size. The results over the claim types show that the default proportion in none of the samples types is significantly different from zero. Furthermore, the default probability for the other 18 claim types is zero. These types are supply, summer-day-camp, service, psychological-care, prescription, pre-school, no claim type, medicine, various, long term care premium, licensed day care adv, learning disability, insurance premium, hospital, home care, hearing or sight impaired aid, hearing, and diagnostic. We do not show them here in the following table to save space. The default proportion for the sample of 8733 claims over all the 25 claim types is 0.001408977, and the same is not significantly different from zero.

Now, since the estimated values of the population default proportion for the 25 claim-types are insignificant, the table puts us into the conjecture that the claims are randomly drawn from population where the default proportion follows Poisson distribution. The findings suggest that the default proportion for notably large sample (say, the population) is 0.001408977. The rest of the study now uses another random sample of 2125 claims.

Precession (98%) Towards Individual Claim:

In order to recognize the pattern of the distribution for sample default proportion, we formulate the following probabilistic method for increased sample size (see, identity (3)), and thereby forward our prediction regarding the population default proportion. Where, P_i is the probability of default issues (an issue may be the place or a cause of default etc.) in an i th claim in the audit process; D_i is the number of default issues in the i th individual claim, and X_i is the number of total issues in that claim.

$$P_i = \frac{D_i}{X_i} \dots \dots \dots (3) \text{ Where, } i = 1, 2, 3, \dots \dots \dots n$$

Now, the probability of default for an individual claim would not be zero, but very small. Thus, if the values of D_i s follow *Poisson* distribution, the distribution of P_i s will be positive but insignificant. To examine this nature of the sample distribution we use to apply z-test. Here, the null hypothesis is that, for a randomly selected sample of the audit claims, the mean value of the values of the P_i s' provides desired 98 % or more level of precession (that is, an error level which is less or equal to 0.02). Here, we hypothetically assume that the population default proportion is not more than 0.02. In our case for the sample size of 2125 insurance claims we find that the estimated population standard error is 0.0030.

The above results show that the mean (referred as "Avr. _ in the table) proportion of the sample default proportions are lower than their expected population default proportion of 0.02 significantly for the case of Line Audit (except for Trainee) and Amount Audit. Here, the same for Claim Audit (except for Sr. Processor) are significantly higher than 0.02 for the both Trainer and Processor. The desired 98% precession level for any sample size therefore cannot be identified towards the audit of individual claims for the Trainee with respect to Line Audit and Claim Audit and for the Processor with respect to the Claim Audit.

Precession (98%) Towards Sample of Claims

The above methodology is now applied for the audit cases on a continuous basis where we require increasing the sample size randomly and we require identifying the probability of default for each of the samples with the increased sizes.

Table 1: Identification of Sample Bias within between Claim Types

Major Claim Types	N	$\sum_{k=1}^N tv_k$	$\sum_{k=1}^N ev_k$	f_l value	variance of f_k	t – value	F – value
VISION	612	256217.07	283	0.00110453	0.001633987	0.00000180	0.996708
TRAVEL EXPENSE	102	7719.35	2.32	0.00030054	0.009803922	0.00000295	
OTC	377	13008.12	106.95	0.00822179	0.005290931	0.00003080	
MEDICAL	3222	1036439.6	2251.51	0.00217235	0.00432759	0.00000252	
LICENSED DAY CARE	889	1654908.2	2599.38	0.00157071	0.001124859	0.00000177	
DENTAL ORTH-ODONTIA	315	179525.58	365	0.00203314	0.009532094	0.00001118	
DENTAL ORDER	1040	558489.86	5472.71	0.00979912	0.006328986	0.00002417	

Notes: All t-values and F-values are not significantly different from zero.

This is shown as below in the identity (4) where, P_j is the probability of default claims for a random sample of size ‘j’ as identified in the audit process, D_j is the number of default issues in the j^{th} size, and X_j is the number of total issues in the j^{th} size of sample.

$$P_j = \frac{\sum_{i=1}^n D_j}{\sum_{i=1}^n X_j} \dots\dots\dots(4), \text{ where, } j = 1, 2, 3, \dots\dots\dots n$$

Since, we have earlier assumed that the claims in the system do enter randomly through the system of data processing; each sample of ‘j’ claims also provides random variables. If the magnitude of D_j follows *Poisson* distribution, then the distribution of P_j for a large sample size of n for any value of j will have positive but insignificant statistics for the mean and variance since the probability of default would be rare but not zero. In order to examine this view regarding the sample distribution we use to apply Z-test. Here, the null hypothesis is the mean value of the values of the P_j s’ provides an error level which is less or equal to the hypothetical population default proportion (i.e., 0.02). In the course of utilizing the values of P_j s however we estimate the population standard error to be 0.0017. The findings for each of the three audit references as audited by Trainee’s, Processor, and Sr.

Processor are shown in Table 3. The results show that the sample proportions are mostly significantly lower than the expected population default proportion of 0.02 for *Amount Audit*. The same for Processor and Sr. Processor with regards to *Line Audit* are significantly lower than expected population default proportion. For the Trainee, the sample default proportion is insignificantly lower than 0.02. The sample proportions in *Claims Audit* are significantly higher than 0.02 specifically for Trainee and Processor. The said proportion for Sr. Processor is insignificantly lower than the population default proportion. Therefore, for the case of *Line Audit* by Trainee, the desired level of precession is not possible to be identified through audit of a sample of claims. Furthermore, the findings suggest that *Claim Audit* is not a suitable audit reference for sample auditing.

Now, since the default proportion (in the above table, size wise) of the sample audit for the audit references of Line Audit and Amount Audit lies below the desired level of 2 percent, the work at hand presently proceeds onto determination of the sample size at 98 percent precessions. In doing so, we would like to estimate the size of the desired sample with 95 % to 99.74% level of confidence. However, with regards to Claims Audit determination of sample size with 98 percent level of precessions is not possible.

Table 2: Identification of Sample Distribution for Default Proportion (claim wise)

		Line Audit		Amount Audit		Claims Audit	
	Parameters	Avr_Sample Proportion	Sample std dev.	Avr_Sample Proportion	Sample std dev.	Avr_Sample Proportion	Sample std dev.
Trainees’ role		0.032405221	0.163095472	0.015529	0.107781	0.080498323	0.272128327
	Z values	4.084660086		-1.4721501		19.92024922	
Processors’ role		0.0109063	0.0835885	0.012524	0.099779	0.02920057	0.16840866
	Z values	-2.994293		-2.46169		3.02946804	
Sr. Processors’ role		0.00739789	0.073557289	0.0070539	0.0760748	0.016924565	0.12902005
	Z values	-4.149489639		-4.2627644		-1.012646842	

Notes: Bold z-values indicate that the z-values are significant 2% level of significance; and “Avr.”_infers mean value.

Determination of Sample Size

The basic understanding of the determination of sample size theoretically considers a statistical set up of normal distribution of the parent population. In the paper, we use to follow a theoretical approximation that once the sample size tends to become large the *Poisson* distribution tends to follow the normal distribution (). Hence, by way of applying this approximation the study forwards a methodology on determination of sample size. The formulas are given as follows. Here, we firstly work out the proportion of the error and then derive the sample size for the finite population size and infinite population size. In these regards we follow the methodology in Kothari (2004, p. 174 - 181).

Estimation of Size for Random Samples: Let us assume that notation p refers the sample proportion where $q = 1 - p$ where z is the value of the standard normal variate at a given confidence level for any sample size n . Then, at the given precision rate, the acceptable error, 'e' can be expressed as follows (see also, Israel, 1992; p. 3).

$$e = z \cdot \sqrt{\frac{p \cdot q}{n}}, \text{ that is, } n = \frac{z^2 \cdot p \cdot q}{e^2}.$$

However, the above formula is designed for sampling with an infinitive population. The same can also be modified for a finite population (N), which is mention as bellow.

$$n = \frac{z^2 \cdot p \cdot q \cdot N}{e^2 (N - 1) + z^2 \cdot p \cdot q}$$

Now, to forward a conservative outlook in determining the sample size, we assume that the desired error will be more than the default proportion. The data in Table: 3 offers that the value of the error can be taken as 0.02 (that is, 98% level of precession) only for *line audit* and *amount audit*. Again, since P_j varies with their respective mean and variance, the study uses to identify each sample size for each of p_j s and q_j s. Thus, the same will provide a different sample size n_j .

Here, the predetermined error value, i.e., the population default proportion is assumed to be 0.02, and thus, the values for p_j s and q_j s would vary with the changes in the sample size n_j . Applying the following equation the values of n_j over our data series are derived:

$$n_j = \frac{z^2 \cdot p_j \cdot q_j}{e^2}, \text{ where } p_j = 1 - q_j; \text{ and}$$

$$n_j = \frac{z^2 \cdot p_j \cdot q_j \cdot N}{e^2 (N - 1) + z^2 \cdot p_j \cdot q_j}, \text{ where } N \text{ is the infinitively}$$

large sample size ($N = 2125$).

We also run the above equation on a continuous basis over all j -s (where $j = 1, 2, 3, \dots, n$) and thereby, derive the magnitudes of n_j for different values of j -s. We report their descriptive statistics in the following heading of results and findings. However, since the sample default proportion follows normal distribution, the best estimate for the desired sample size would be the mean sample size, which is also an unbiased estimator.

RESULTS AND FINDINGS

The results and findings are given in the following tables. Here, the mean statistics along with their standard deviation refer the sample size.

- (i) In the case of the said audit by Trainee, in Table: 4 for auditing with reference to a finite population size of 2125, the findings suggest that in order to obtain the required 98% level of precession, the confidence level of 95% can be observed for a sample audit with size of 158 claims (with a std. dev. of 58 claims) for line audit, 122 claims (with a std. dev 58 claims) for amount audit, and 547 claims (with a std. dev. of 178 claims) without any reference to the amount and the lines in each claim for the number of the claim audit.

Table 3 : Identification of Sample Distribution for Default Proportion (size wise)

	Parameters	Line Audit		Amount Audit		Claims Audit	
		Avr. Sample Proportion	Sample std dev.	Avr. Sample Proportion	Sample std dev.	Avr. Sample Proportion	Sample std dev.
Trainees' role		0.017269563	0.008040481	0.013137418	0.007226847	0.083894	0.039402
	Z values	-0.899049488		-2.259638659		21.03831	
Processors' role		0.009285561	0.001752426	0.013025058	0.005506838	0.037516	0.007492
	Z values	-3.527937269		-2.296635393		5.767553	
Sr. Processors' role		0.007226618	0.007335451	0.004469385	0.002597399	0.018799	0.02719
	Z values	-4.205884384		-5.113757018		-0.3953	

Notes: Bold z-values indicate that the z-values are significant 2% level of significance and "Avr."_infers mean value.

Table 4: Sample size (Claims to be audited) For Trainee¹

<i>Cases</i>	<i>Finite Population</i>	<i>infinite Population</i>					
		<i>Level of confidence</i>					
	95 %	95 %	96 %	97 %	98	99 %	99.74 %
Line Audit							
Mean	157.9267	172.4545	197.9628	217.2687	258.5803	303.4526	403.9697
Std. dev	57.59424	65.21771	74.8426	82.136	97.74673	114.7158	152.7317
Max Value	241	272	312	342	407	478	636
Min. Value	13	13	15	16	19	23	30
Mode Value	180	197	226	248	267	313	417
Amount Audit							
Mean	122.1939	131.3939	150.8473	165.5779	197.0396	231.2527	307.8426
Std. dev	57.28945	62.95252	72.25448	79.31146	94.36958	110.7428	147.4425
Max Value	207	229	263	288	343	403	536
Min. Value	8	8	9	10	12	14	19
Mode Value	141	151	173	180	226	265	353
No. Claims Audit							
Mean	547.1308	768.1196	881.7598	967.3796	1151.238	1351.114	1798.771
Std. dev	178.0824	282.6068	324.4231	356.3502	424.1106	497.7495	662.7606
Max Value	755	1171	1344	1475	1755	2060	2742
Min. Value	58	60	68	75	89	105	140
Mode Value	618	872	1001	1098	1304	1532	2047

Table 5: Sample size (Claims to be audited) For Processor²

<i>Cases</i>	<i>Finite Population</i>	<i>infinite Population</i>					
		<i>Level of confidence</i>					
	95 %	95 %	96 %	97 %	98	99 %	99.74 %
Line Audit							
Mean	85.8559	89.51771	102.7676	112.7977	134.2149	157.5308	209.7239
Std. dev	12.1516	13.09915	15.03836	16.50502	19.6364	23.03877	30.68548
Max Value	113	119	137	151	179	210	280
Min. Value	35	36	41	45	53	62	83
Mode Value	77	80	92	101	120	141	188
Amount Audit							
Mean	120.0363	128.1069	146.9547	161.3254	192.01	225.3274	300.0144
Std. dev	40.00448	45.97622	52.86458	57.92884	68.90521	80.90576	107.695
Max Value	252	285	328	360	428	502	669
Min. Value	75	78	0.037516	98	117	137	182
Mode Value	79	82	94	103	123	144	192

1 Even though, we like to propose about sample size for auditing for the case of number of claims for Trainee, it should be used cautiously, since we our population proportion of error is 0.02 (assumed) and our sample proportion is 0.083894. These sample size will not be as per out desired outputs.

2 Even though, we like to propose about sample size for Auditing for the case of Number of Claims for Processor, it should be used cautiously, since we our population proportion of error is 0.02 (assumed) and our sample proportion is 0.037516. The sample size will not produce our desired outputs.

<i>Cases</i>	<i>Finite Population</i>	<i>infinite Population</i>					
		<i>Level of confidence</i>					
	95 %	95 %	96 %	97 %	98	99 %	99.74 %
No. Claims Audit							
Mean	300.4682	350.9534	402.8787	442.1669	526.2159	617.5687	822.1975
Std. dev	39.44509	54.12834	62.1417	68.19136	81.14763	95.23882	126.7972
Max Value	400	493	566	621	739	867	1154
Min. Value	196	216	248	272	323	380	505
Mode Value	266	304	348	383	455	535	702

Table 6: Sample size (Claims to be audited) For Sr. Processor

<i>Cases</i>	<i>Finite Population</i>	<i>infinite Population</i>					
		<i>Level of confidence</i>					
	95 %	95 %	96 %	97 %	98	99 %	99.74 %
Line Audit							
Mean	65.13491	68.3912	78.49807	86.15329	102.5358	120.3409	160.2186
Std. dev	42.05966	60.65686	69.62109	76.41476	90.9465	106.7239	142.0938
Max Value	892	1537	1764	1936	2304	2704	3600
Min. Value	21	21	24	26	31	37	49
Mode Value	53	55	63	69	82	96	128
Amount Audit							
Mean	42.2105	43.30471	49.7105	54.54956	64.91462	76.17468	101.4303
Std. dev	21.14309	23.83901	27.3724	30.04074	35.73549	41.95353	55.86086
Max Value	317	372	427	469	558	655	872
Min. Value	15	15	18	19	23	27	36
Mode Value	39	40	46	50	60	70	93
No. Claims Audit							
Mean	154.5438	170.1447	195.3145	214.3633	255.1132	299.3957	398.5999
Std. dev	66.74968	108.695	124.7665	136.9435	162.9697	191.2674	254.6439
Max Value	1128	2401	2756	3025	3600	4225	5625
Min. Value	43	43	50	55	65	76	102
Mode Value	154	166	191	210	250	293	391

Table 7: Sample size (Claims to be audited) For Trainee³

<i>Cases</i>	<i>At 95 % Level of confidence With 20 % decline</i>						
	<i>Population Range</i>	<i>Bellow 2125</i>	<i>2125-5000</i>	<i>5000-10000</i>	<i>10000-20000</i>	<i>20000-50000</i>	<i>50000-100000</i>
Line Audit	Sample Size	160	320	544	924.8	1965.2	3340.84
	Percentage	7.5294	6.4	5.44	4.624	3.9304	3.34084
Amount Audit	Sample Size	125	250	425	722.5	1535.3	2610.031
	Percentage	5.882	5	4.25	3.6125	3.070	2.610031
No. Claims Audit	Sample Size	547	1094.26	1860.24	3162.41	6720	11424.23
	Percentage	25.747	21.885	18.6024	15.812	13.44	11.42423

³ Even though, we like to propose about sample size for Auditing for the case of Number of Claims for Trainee, it should be used cautiously, since we our population proportion of error is 0.02 (assumed) and our sample proportion is 0.083894. These sample size will not desired outputs.

Table 8: Sample size (Claims to be audited) For Processor⁴

Cases	At 95 % Level of confidence With 20 % decline						
	Population Range	Bellow 2125	2125-5000	5000-10000	10000-20000	20000-50000	50000-100000
Line Audit	Sample Size	85	160	256	409.6	819.2	1310.72
	Percentage	4	3.2	2.56	2.048	1.6384	1.31072
Amount Audit	Sample Size	120	225.88	361.41	578.25	1156.52	1850.43
	Percentage	5.647	4.5176	3.6141	2.8912	2.313	1.850
No. Claims Audit	Sample Size	300	564.70	903.52	1445.64	2891.29	4626.07
	Percentage	14.11	11.294	9.0352	7.228	5.7825	4.62

Table 9: Sample size (Claims to be audited) For Sr. Processor

Cases	At 95 % Level of confidence With 20 % decline						
	Population Range	Bellow 2125	2125-5000	5000-10000	10000-20000	20000-50000	50000-100000
Line Audit	Sample Size	65	122.35	195.764	313.223	626.44	1002.315
	Percentage	3.058	2.44	1.957	1.566	1.252	1.002
Amount Audit	Sample Size	42	79.05	126.49	202.39	404.78	647.64
	Percentage	1.976	1.58	1.2649	1.012	0.809	0.64
No. Claims Audit	Sample Size	154	289.88	463.811	742.09	1484.19	2374.7
	Percentage	7.247	5.797	4.638	3.711	2.96	2.37

⁴ Even though, we like to propose about sample size for Auditing for the case of Number of Claims for Processor, it should be used cautiously, since we our population proportion of error is 0.02 (assumed) and our sample proportion is 0.037516. These sample size will not desired outputs.

(ii) In the case of audit by Processor, in Table: 5 with reference to finite population, the findings suggest that in order to obtain required 98% level of precession with the confidence level of 95% the same could be observed through a sample audit with a sample size of 86 claims (with a std. dev. of 12 claims) for *line audit*, 120 claims (with a std. dev. of 40 claims) for *amount audit*, and 300 claims (with a std. dev. of 39 claims) with out any reference to the amount and the lines in each claim for number of *claims audit*.

(iii) In the case of Sr. Processor, in Table: 6, the findings suggest that in order to obtain the required 98% level of precession with confidence level of 95% for a finite population, one can audit the sample claims through a sample audit with the sample size of 65 claims (with a std. dev. of 42 claims) for *line audit*, 42 claims (with a std. dev. of 21 claims) for *amount audit*, and 155 claims (with a std. dev. of 67 claims) with out any reference to the amount and the lines in each claim for number of *claims audit*.

In all the above cases respectively for *line audit*, *amount audit*, and number of *claims audit* for an infinite population,

the sample size at 95 % level of confidence becomes (i) 172 claims, 131 claims, and 768 claims for an audit by Trainee; (ii) 90 claims, 128 claims, and 351 claims for the same by Processor; and (iii) 68 claims, 43 claims, and 170 claims for auditing by Sr. Processor. In all the cases for the Trainee, Processor, and Sr. Processor with respect to an infinite size of population, the sample size increases with a higher level of confidence.

SAMPLE SIZE ESTIMATION FOR DIFFERENT RANGES

In the earlier heading, we observe that at different sample size drawn from the population, the default proportion varies. Now, we assume that in an unbiased and random sampling, sample default proportion will be downward sloping and convex to the origin for an increased sample size. In order to extend the above observation for a hypothetical population size we now forward that the default proportion will fall but at a slower rate. However, we like to propose the following as our recommendation. We suggest the decline in the sample size would be at a diminishing rate of 15% for Trainee, 20% for Processor, and 25% for Sr. Processor. The logical explanation for this diversity is that (see also Table 3) a lesser default proportion is related firstly with Sr. Processor and then to Processor than that of trainee. However, we take the percentage figures on an ad-hoc basis.

CONCLUSION

The present work has tried to examine the determination of sample size for auditing of medical insurance claims by an audit firm in its business to the offshore clients. Even if sample size determination is a tougher task for an auditor specifically at times of huge fluctuations with in the population parameters, the study examines the nature of the medical insurance claims and the default rates those involved with the different claim types and under audit by the different types / layers of the auditors. The study explores that the distribution of the default claims is mostly *Poisson* in nature and our binomial extension supports that the sample size for the Processor and the Sr. Processor is easier to determine to ensure 98% level of precession with a 95% level of confidence. However, in case of an aggregate audit of the claims by the trainee care should be taken. The study also suggests sample sizes for different population sizes. The opinions provided in the study however follow a conservative outlook with respect to the sample size for different ranges of population size. We observe that the percentage of sample size to be audited will move at a diminishing rate as the population size increases.

REFERENCES

- Adcock, C.J., (1992). Bayesian Approaches to the Determination of Sample Sizes for Binomial and Multinomial Sampling-Some Comments on the Paper by Pham-Gia and Turkkan, *The Statistician* , 41 (4), 399-404
- Adcock, C.J. (1997). Sample size determination: A review. *The Statistician*, 46 (2), 261–283
- Gelman, A., (2008). Objections to Bayesian statistics, *Bayesian Analysis*, 3(3), 445-450, DOI:10.1214/08-BA318
- Israle, G. D., (1992). Determining Sample Size, Fact Sheet PEOD - 6, November, 1992; Florida Co-operative Extension Service, University of Florida.
http://www.soc.uoc.gr/socmedia/papageo/metaptyxiakoi/sample_size/samplesize1.pdf
- Kothari, C. R., (2004). *Research Methodology, Methods & Techniques*, New Age International (P) Limited, Publishers (formerly Willy Eastern Limited), New Delhi, 2nd Edition.
- Kruschke, J. K., (2010). What to believe: Bayesian methods for data analysis. *Trends in Cognitive Science*, 14 (7), 293–300
- Pathak, J., Vidyarthi, N., & Summers, S. L., (2005). A fuzzy-based algorithm for auditors to detect elements of fraud in settled insurance claims. *Managerial Auditing Journal*, 20 (6), 632 – 644