

COMPARISON OF RANDOM, SYSTEMATIC, SIEVE, CELL, AND LAHIRI SAMPLING METHODS IN MONETARY UNIT SAMPLING: A CASE STUDY OF LOCAL TAX RECEIVABLES

Aan Subarkah^{*}, Georgina Maria Tinungki^{**}, Nurtiti Sunusi^{**}

^{*} Master of Statistics Program, Faculty of Mathematics and Natural Sciences, Hasanuddin University

^{**} Department of Statistics, Faculty of Mathematics and Natural Sciences, Hasanuddin University

Abstract- This study evaluates the performance of five sampling selection methods in Monetary Unit Sampling (MUS) applied to local tax receivables data from Maros Regency Government in 2023. The research compares Random, Systematic, Cell, Lahiri, and Sieve sampling methods based on their beta risk and upper bound estimation accuracy. Using simulation with 110,616 iterations across 222 scenarios, the study examines seven types of local taxes with a total population value of IDR 43,984,236,590. Results indicate that all methods achieved 0% beta risk, successfully detecting 100% of material misstatements (5% overstatement). However, Sieve Sampling demonstrated superior efficiency with an Efficiency Index 4.59 times higher than other methods, utilizing 65% smaller sample sizes while maintaining detection effectiveness. The findings provide empirical evidence for audit practitioners in selecting optimal sampling methods for public sector financial audits in Indonesia.

Index Terms- Beta Risk, Local tax receivables, Monetary unit sampling, Public sector audit, Sampling methods.

I. INTRODUCTION

Auditing represents a systematic process of obtaining and evaluating objective evidence regarding economic activities and events to establish the degree of correspondence between assertions and established criteria (Mulyadi, 2010). In conducting examinations, auditors employ sampling techniques rather than examining entire populations, aiming to provide an adequate basis for drawing conclusions about the population (Institut Akuntan Publik Indonesia, 2021; Sukrisno Agoes, 2004). While emerging technologies increasingly enable full population testing of client-internal data (Li, Brazel, & Gold, 2024), sampling methods remain essential for obtaining sufficient appropriate audit evidence, particularly when incorporating external evidence sources.

Monetary Unit Sampling (MUS) has emerged as a prevalent statistical sampling method in auditing, implementing Probability Proportional to Size (PPS) where each monetary unit serves as the sampling unit. This approach inherently assigns higher selection probability to transactions with larger values, aligning with auditors' focus on material items (Carrizosa, 2012; Gillett, 2000; Guy et al., 1998; Higgins & Nandram, 2009). While artificial intelligence is increasingly transforming audit processes

through automation and data analytics (Ghafar et al., 2024), traditional sampling methodologies remain fundamental to obtaining sufficient appropriate audit evidence, particularly in contexts requiring professional judgment and external evidence validation.

Various selection methods exist within MUS, including simple random sampling, systematic sampling, cell sampling, sieve sampling, and Lahiri sampling (Horgan, 1994, 1998; Wurst et al., 1989a). Previous research has examined performance differences among these methods. Wurst (1989) demonstrated that cell and sieve sampling produce more accurate upper bound estimates than random sampling for non-small sample sizes. Horgan (1996) confirmed the superiority of cell and sieve sampling over random sampling using Moment Bound evaluation.

However, existing research primarily utilized accounting data in US dollars with relatively small transaction values compared to Indonesian financial contexts. Wurst (1989) examined transactions ranging from \$0.10 to \$24,928.60, while Hoogduin (2015) studied values between \$1 and \$5. Limited research exists on MUS method comparison using Indonesian public sector financial data, creating a gap this study addresses.

This research aims to: (1) test whether beta risks produced by Random, Systematic, Lahiri, Sieve, and Cell sampling methods differ significantly, and (2) identify the most reliable sampling selection method for MUS application in Indonesian public sector auditing.

II. THEORETICAL FRAMEWORK

A. Monetary Unit Sampling

MUS represents an application of Probability Proportional to Size (PPS) sampling where selection probability is proportional to item value (Latpate et al., 2021; Skinner, 2014). In financial auditing, MUS enables efficient estimation of total population misstatement by emphasizing larger-value items with higher misstatement potential (Higgins & Nandram, 2009).

The sample size calculation under MUS employs the hypergeometric distribution, representing the most accurate approach for sampling without replacement scenarios common in auditing (Hoogduin et al., 2010; Stewart, 2012). The sample size (m) is determined as the smallest integer satisfying:

$$\sum_k \frac{\binom{T-E}{m-k} \binom{E}{k}}{\binom{T}{m}} \leq \beta$$

Where:

m = sample size (the smallest integer satisfying the inequality)

k = number of allowable errors

E = total implied error at tolerable error rate

β = planned beta risk

T = total book value

B. Sampling Selection Methods

Random Selection, frequently termed Unrestricted Random Sampling (URS) or Simple Random Sampling of Monetary Units (SRS), is a widely accepted procedure within auditing (Horgan, 1994). Random Selection involves generating random numbers to select monetary units, ensuring each unit maintains equal selection probability. The method requires cumulative totaling and random number sorting (Wurst et al., 1989a). While this traditional approach maintains rigorous statistical validity, it presents certain implementation challenges that must be carefully considered. The procedure necessitates the pre-calculation of cumulative book value subtotals across the entire population and subsequent sorting of generated random numbers to identify the corresponding line items, making it computationally intensive and potentially time-consuming, particularly when dealing with large accounting populations (Horgan, 1998).

Systematic Selection divides the population into equal intervals, selecting units systematically from a random starting point. While offering consistent sample sizes, it may produce unreliable beta risk estimates when errors exhibit periodic patterns (Hoogduin et al., 2015). The fundamental advantage of systematic selection lies in its operational efficiency and predictable outcome characteristics. The method consistently delivers the exact target sample size of distinct line items, eliminating uncertainty about actual sample composition (Horgan, 1998). It requires significantly less computational effort compared to unrestricted random sampling, as it avoids the need for extensive random number generation and sorting procedures (AICPA, 2016; Horgan, 1998). Systematic selection demonstrates superior statistical precision compared to random sampling (Horgan, 1994; Horgan, 1998), particularly in populations containing large line items where substantial gains in precision occur with larger sample sizes. The method also provides automatic coverage of line items exceeding predetermined monetary thresholds, ensuring comprehensive examination of high-value items (AICPA, 2016).

However, systematic selection suffers from critical methodological limitations that compromise its reliability in audit applications. The primary weakness involves its vulnerability to periodic error patterns within accounting populations, which can lead to severely biased sample selection and unreliable risk assessments (Hoogduin et al., 2015).

Sieve Sampling eliminates the need for cumulative totaling by independently evaluating each transaction against a random threshold. Transactions exceeding the sampling interval are

automatically selected, while others undergo probability-based selection (Horgan, 1998). This method represents a list-sequential approach that exploits the natural line item structure of accounting populations, fundamentally distinguishing it from traditional draw-sequential methods that require extensive cumulative calculations (Wurst et al., 1989). The technique operates by generating a random determination for each line item in the population to indicate whether it should be included in the sample and, if selected, which specific monetary unit within that line item represents the sample selection. Research demonstrates that sieve sampling produces unbiased point estimators with superior precision compared to simple random sampling of monetary units, particularly when sample sizes are not excessively small, suggesting potential advantages for bounds-based evaluation methods in audit applications (Wurst et al., 1989). However, sieve sampling exhibits a fundamental characteristic that distinguishes it from fixed sample size methods: inherent variability in achieved sample sizes. Empirical evidence from Horgan (1998) demonstrates that for a target sample size of 100, sieve sampling can yield between 74-129 distinct line items (mean = 99.78, standard deviation = 9.07), representing substantial deviation from predetermined targets. This variability is not a methodological flaw but rather an inherent property of the Poisson sampling category, which includes sieve sampling as proposed by Hájek (1964) to address replacement issues while accepting sample size uncertainty (Horgan, 1998). The method belongs to a class of variable sample size approaches alongside unrestricted random and Lahiri sampling, contrasting with fixed sample size methods such as systematic and cell sampling that consistently achieve target sample sizes.

The primary limitation of sieve sampling involves its inherently variable achieved sample size, which may deviate substantially from the predetermined target depending on the randomly generated thresholds and underlying population characteristics (Horgan, 1998). This variability creates significant planning uncertainties, as the actual number of line items selected can only be predicted within certain probabilistic limits rather than precisely determined at the design stage.

Cell Sampling modifies systematic selection by performing independent selection within each interval (cell), addressing systematic sampling's bias risks but potentially resulting in multiple selections of the same item (Leslie et al., 1979). This method represents a stratified sampling approach where the population of monetary units is divided into n cells of equal size $T(Y)/n$, with one monetary unit selected randomly from each cell independent of selections in other cells (Wurst et al., 1989). The technique eliminates the systematic bias vulnerability associated with periodic error patterns by introducing random selection within each stratum, while maintaining the advantageous property of ensuring sample distribution across the entire population. Cell sampling preserves the efficiency benefits of systematic selection while providing enhanced statistical rigor through its independence assumption, making it a popular compromise between the precision of systematic sampling and the theoretical soundness of random selection methods.

The primary limitation of cell sampling lies in its potential for multiple selections from the same line item when large items

straddle multiple cells, potentially reducing the number of distinct line items below the target sample size (Horgan, 1998). This "multiple hits" phenomenon may compromise the effectiveness of audit coverage, as fewer unique line items receive examination than originally planned.

Lahiri Sampling enables selection before complete population availability, utilizing natural transaction ordering but potentially resulting in higher selection variance compared to other methods (Horgan, 1994). This method operates by first selecting a random number between one and the total number of line items to identify a candidate for consideration, followed by generating a second random number between one and the maximum book value (or an estimated maximum based on auditor experience) to determine whether the identified line item should be included in the sample (Horgan, 1998). The selection process continues iteratively until the required sample size is obtained, making it particularly suitable for situations where the auditor wishes to commence sampling procedures before complete population enumeration is available.

Lahiri sampling offers exceptional implementation flexibility by enabling sample selection to commence before complete population availability, making it particularly valuable in time-constrained audit environments where early testing is advantageous. The method eliminates the computational burden associated with cumulative subtotal calculations, requiring only basic arithmetic operations that can be performed manually without sophisticated computer assistance. The primary limitation of Lahiri sampling lies in its consistently inferior statistical precision compared to alternative monetary unit sampling methods, demonstrating performance similar to unrestricted random sampling with consistently higher variability in bound estimates (Horgan, 1998).

C. Evaluation Procedures

Cell Bound evaluation procedures, developed by Leslie et al. (1979), calculate upper misstatement bounds by combining statistical distribution-based error rate bounds with observed taint percentages. This method provides reliable coverage probability while avoiding excessive conservatism (Bimpeh, 2006; Hoogduin et al., 2010). Recent advances in bound estimation include weighted empirical likelihood approaches that address the conservative nature of traditional bounds like the Stringer bound, particularly suitable for skewed error distributions common in MUS applications (Berger et al., 2021). Additionally, Bayesian approaches to MUS evaluation have gained attention, with Stewart and Sunderland (2024) demonstrating how gamma distributions can model audit assurance profiles and integrate prior assurance with sampling evidence through Bayes' rule, offering a theoretically rigorous framework for updating confidence levels based on sample results.

Beta risk, representing the risk of incorrect acceptance, constitutes auditors' primary concern as failure to detect material misstatements compromises audit effectiveness and may result in legal consequences (Whittington et al., 2010; Woodhead, 1997).

III. RESEARCH METHODOLOGY

A. Data Source and Variables

This study utilizes primary data from Maros Regency Government's 2023 financial statements, specifically local tax receivables recorded in the financial reports. The data represents actual tax receivables from seven tax types: Groundwater Tax, Entertainment Tax, Hotel Tax, Non-Metal Mineral and Rock Tax, Advertisement Tax, Restaurant Tax, and Swallow's Nest Tax. The complete dataset comprises 304 transactions with a total population value of IDR 43,984,236,590.

The variables employed in this research include:

- **Tax Type:** Categorical variable identifying one of seven local tax categories, used to stratify the population into distinct audit populations.
- **Transaction ID:** Unique identifier for each tax receivable transaction, serving as the primary reference for sample selection and traceability.
- **Receivables Value:** The monetary amount of each tax receivable transaction, representing the book value used in Monetary Unit Sampling calculations.

Table 1 presents the data structure used in this study.

Table 1. Data Structure of Maros Regency Tax Receivables 2023

No	Tax Type	Transaction ID	Receivables (IDR)
1	Groundwater	ABT1	3,951,000.00
2	Groundwater	ABT2	1,778,480.00
...	...	...	...
304	Swallow's Nest	WAL4	1,000,000.00

Source: Maros Regency Government Financial Statements Fiscal Year 2023

B. Simulation Design

The simulation design follows a methodology similar to Hoogduin et al. (2015), adapted for Indonesian public sector financial data. The framework incorporates:

- Error rates: 5% of total population value (tolerable error threshold)
- Overstatement levels: 50%, 60%, 70%, 80%, 90%, and 100%
- Allowable errors (k): 0 and 1
- Beta risk levels (β): 0.01, 0.05, 0.15, and 0.29
- Iterations: 100 per parameter combination

The combination of parameters resulted in 222 distinct study populations for analysis. This total derives from the successful sample size calculations (those not exceeding population size) across all seven tax types, multiplied by six overstatement levels. Table 2 presents the distribution of study populations by tax type.

Table 2. Distribution of Study Populations by Tax Type

Tax Type	Valid (k, β) Combinations	Overstatement Levels	Study Populations
Groundwater	4	6	24
Entertainment	5	6	30
Hotel	4	6	24
Non-Metal Mineral	3	6	18
Advertisement	8	6	48
Restaurant	8	6	48
Swallow's Nest	5	6	30
Total	37	6	222

Each of the 222 study populations was subjected to 100 iterations across all five sampling methods, generating 110,616 total sampling iterations: Random Sampling (22,200 iterations), Systematic Sampling (22,200 iterations), Cell Sampling (22,200 iterations), Lahiri Sampling (22,200 iterations), and Sieve Sampling (21,816 iterations). The slight variation in Sieve Sampling iteration counts reflects its inherent variable sample size characteristic.

Sample sizes were calculated using the hypergeometric distribution, consistent with widely accepted audit sampling practice (Hoogduin et al., 2010, 2015; Stewart, 2012). Evaluation was performed using Cell Bound procedures, known for reliability in MUS applications (Leslie et al., 1979; Hoogduin et al., 2015). Material overstatements were randomly inserted into populations, and sampling methods were applied to detect these errors. This approach enables systematic comparison of sampling method effectiveness under controlled conditions while maintaining consistency with established audit research methodology.

IV. RESULTS AND DISCUSSION

A. Descriptive Analysis

The research data exhibits significant heterogeneity across seven tax types comprising 304 transactions with a total population value of IDR 43,984,236,590. Tables 3 and 4 present the descriptive statistics of the tax receivables data.

Table 3. Population Characteristics of Tax Receivables Data

Tax Type	Total Population (IDR)	Number of Transactions	Min (IDR)	Max (IDR)
Groundwater	56,364,740	20	40	12,780,000
Entertainment	1,316,700	3	100,000	866,700
Hotel	49,120,501	15	1	12,780,000
Non-Metal Mineral	40,288,516,531	24	725	29,730,628,174
Advertisement	372,351,631	71	100	38,855,000
Restaurant	3,213,066,487	167	1	1,650,846,906
Swallow's Nest	3,500,000	4	500,000	1,200,000
Total	43,984,236,590	304	1	29,730,628,174

Table 4. Central Tendency and Dispersion Measures

Tax Type	Mean (IDR)	Coefficient of Variation (%)
Groundwater	2,818,237	119.02
Entertainment	438,900	89.09
Hotel	3,274,700	109.50
Non-Metal Mineral	1,678,688,189	360.64
Advertisement	5,244,389	142.60
Restaurant	19,239,919	682.71
Swallow's Nest	875,000	34.13
Total	144,684,989	

Based on Table 3, Non-Metal Mineral and Rock Tax dominates the population value, representing IDR 40,288,516,531 (91.60% of total), followed by Restaurant Tax at IDR 3,213,066,487 (7.31%). The smallest population value belongs to Entertainment Tax at IDR 1,316,700 (0.003%).

In terms of transaction frequency, Restaurant Tax exhibits the highest count with 167 transactions (54.93% of total), while Entertainment Tax has the lowest with only 3 transactions (0.99%). The population demonstrates extreme value heterogeneity, with transaction values ranging from IDR 1 to IDR 29,730,628,174, yielding a maximum-minimum ratio of 29,730,628,174:1. This finding aligns with Hoogduin et al. (2015), who observed that audit populations frequently exhibit highly uneven distributions with certain items having substantially larger values than others.

As shown in Table 4, most tax types display coefficient of variation exceeding 100%, indicating highly skewed distributions characteristic of audit populations. Restaurant Tax shows the highest variability (CV=682.71%), while Swallow's Nest Tax demonstrates the most homogeneous distribution (CV=34.13%). According to Hoogduin et al. (2015), populations with such high coefficients of variation present particular challenges for sampling methods, especially regarding consistency in projected error estimation.

Based on these population characteristics showing high heterogeneity and skewed distributions, sample sizes were determined for each tax type across different risk tolerance levels using hypergeometric distribution principles.

B. Sample Size Determination

Sample size determination follows the hypergeometric distribution approach established by Hoogduin et al. (2015), with tolerable error set at 5% of population value and allowable errors (k) of 0 and 1. This method provides more accurate results for sampling without replacement compared to Poisson or binomial approximations, particularly suitable for audit applications where replacement is not permitted.

Sample size calculations reveal variation across tax types and risk levels. For $\beta=0.01$ and $\beta=0.05$, several tax types require sample sizes exceeding population size, particularly for smaller populations. Table 5 presents sample sizes for k=0 scenario.

Table 5. Sample Sizes for $k=0$ (No Allowable Errors)

Tax Type	$\beta=0.01$	$\beta=0.05$	$\beta=0.15$	$\beta=0.29$
Groundwater	NA	NA	37	25
Entertainment	NA	59	37	25
Hotel	NA	NA	37	25
Non-Metal Mineral	NA	NA	NA	25
Advertisement	42	42	37	25
Restaurant	37	37	37	25
Swallow's Nest	NA	59	37	25

Note: NA indicates required sample size exceeds population size.

Using these calculated sample sizes, five sampling methods were applied across all 222 study scenarios to evaluate their comparative performance in detecting the 5% material overstatements inserted into each population.

C. Sampling Method Performance

All five sampling methods achieved 0% beta risk across 110,616 iterations, successfully detecting all material misstatements. However, significant differences emerged in efficiency metrics:

Table 6. Comparative Performance Metrics

Method	Beta Risk (%)	Mean Projected Error (%)	Sample Size	Efficiency Index
Random	0	66.98	37 (fixed)	1.00
Systematic	0	66.92	37 (fixed)	1.06
Cell	0	66.93	37 (fixed)	1.13
Lahiri	0	66.90	37 (fixed)	1.17
Sieve	0	98.74	11-14 (variable)	4.59

Table 6 demonstrates that while all methods achieved identical beta risk (0%), substantial differences exist in efficiency metrics. The Efficiency Index, which considers detection rate, accuracy, sample size, and variability, reveals significant performance disparities. Sieve Sampling achieves the highest mean projected error (98.74%), indicating superior accuracy as this value approximates the actual 100% overstatement level more closely than other methods (~66.9%). This enhanced accuracy is achieved with substantially smaller sample sizes (11-14 units) compared to fixed-size methods (37 units), resulting in an Efficiency Index 4.59 times higher than Random Sampling. The fixed-size methods (Random, Systematic, Cell, Lahiri) show similar mean projected errors but differ slightly in efficiency indices due to their inherent methodological characteristics.

D. Efficiency Analysis

Sieve Sampling demonstrated superior efficiency with an index 4.59 times higher than Random Sampling. Despite utilizing 65% smaller sample sizes (mean=13 versus 37), Sieve Sampling maintained 100% detection effectiveness while producing projected errors closest to actual misstatement levels (98.74% versus actual 100%).

The variable sample size characteristic of Sieve Sampling (11-14 units) presents planning challenges but offers substantial efficiency gains. These findings align with established empirical

evidence from Horgan (1998), who documented that sieve sampling consistently produces variable sample sizes with substantial standard deviations, yet maintains superior statistical precision compared to traditional methods. The observed sample size range of 11-14 units in this study corresponds to the theoretical expectation of variability documented in the literature, where actual sample sizes can deviate significantly from predetermined targets while preserving detection effectiveness.

Systematic Sampling showed the highest consistency in error detection (5-7 errors per iteration), while Lahiri Sampling exhibited the widest detection range (2-10 errors). These results support the methodological classification distinguishing fixed sample size methods (Systematic, Cell) from variable sample size methods (Sieve, Random, Lahiri), as theoretical research suggests.

These efficiency findings, demonstrating Sieve Sampling's superior performance alongside the reliable detection capability of all methods, have important practical implications for auditors in the Indonesian public sector.

V. IMPLICATIONS AND LIMITATIONS

This study provides empirical evidence supporting Sieve Sampling's efficiency advantage in Indonesian public sector auditing contexts. The 65% reduction in sample size without compromising detection effectiveness offers significant practical benefits for audit resource optimization.

The findings contribute to the theoretical understanding of MUS method classification by empirically validating the distinction between fixed and variable sample size approaches. Fixed sample size methods (Systematic, Cell) provide predictable resource requirements and consistent coverage, making them suitable for environments requiring precise audit planning. Conversely, variable sample size methods (Sieve, Random, Lahiri) offer methodological advantages such as computational efficiency and bias reduction, albeit with planning uncertainties. Sieve sampling emerges as the most efficient within the variable category, achieving superior precision with significantly reduced sample sizes despite inherent size variability.

The empirical evidence aligns with Horgan's (1998) theoretical framework indicating that sample size variability represents a fundamental methodological trade-off rather than a limitation. While variable sample sizes complicate audit planning, they enable computational simplicity and enhanced statistical precision, particularly valuable in resource-constrained public sector environments where efficiency gains can substantially impact audit feasibility and cost-effectiveness.

However, several limitations warrant consideration. The study examined only overstatement errors, excluding understatement scenarios. Additionally, the single-period, single-location data limits generalizability. While technological advances increasingly enable full population testing of client-internal data (Li et al., 2024), this research demonstrates the continued relevance of sampling methods for audit quality, particularly when appropriateness of evidence sources remains a critical consideration. Future research should incorporate multiple error types, extended time periods, and diverse geographical contexts.

Furthermore, while this study employed Cell Bound evaluation procedures, recent developments in weighted empirical likelihood approaches offer promising alternatives that may provide less conservative bounds with better coverage properties for skewed error distributions (Berger et al., 2021). Future research could also explore hierarchical Bayesian approaches for improving parameter estimation precision in MUS contexts by leveraging cross-sectional information across multiple audit entities (Chintha & Kallapur, 2024). Additionally, investigating how AI-driven technologies such as machine learning and natural language processing (Ghfar et al., 2024) can complement traditional MUS methodologies represents a promising avenue for enhancing both sampling efficiency and audit quality in public sector contexts.

VI. CONCLUSION

This research demonstrates that while all five MUS sampling methods effectively detect material misstatements with 0% beta risk, substantial efficiency differences exist. Sieve Sampling emerges as the most efficient method, achieving 4.59 times higher efficiency than traditional approaches through significantly reduced sample sizes without sacrificing detection capability.

These findings contribute to audit methodology literature by providing empirical evidence from Indonesian public sector data, addressing the gap in non-Western, large-value transaction contexts. Auditors can leverage these insights to optimize sampling strategies, potentially reducing audit workload by up to 65% while maintaining effectiveness.

Future research should explore combined overstatement-understatement scenarios, compare alternative bound calculation methods, and validate findings across diverse public sector entities and time periods.

REFERENCES

- [1] American Institute of Certified Public Accountants. (2016). Audit guide: Audit sampling. John Wiley & Sons.
- [2] Berger, Y. G., Chiodini, P. M., & Zenga, M. (2021). Bounds for monetary-unit sampling in auditing: an adjusted empirical likelihood approach. *Journal of Statistical Research*, 45(3), 1-25.
- [3] Bimpeh, Y. (2006). Statistical Modelling and Inference for Financial Auditing. *International Journal of Auditing*, 10(2), 125-142.
- [4] Carrizosa, E. (2012). On approximate Monetary Unit Sampling. *European Journal of Operational Research*, 217(2), 479-482.
- [5] Chintha, B. R., & Kallapur, S. (2024). Hierarchical Bayesian Models in Accounting Research. *Contemporary Accounting Research*, 41(2), 285-315.
- [6] Ghfar, I., Perwitasari, W., & Kurnia, R. (2024). The Role of Artificial Intelligence in Enhancing Global Internal Audit Efficiency: An Analysis. *Asian Journal of Logistics Management*, 3(2), 64-89.

- [7] Gillett, P. R. (2000). Monetary unit sampling: a belief-function implementation for audit and accounting applications. *International Journal of Approximate Reasoning*, 25(1), 43-70.
- [8] Higgins, H. N., & Nandram, B. (2009). Monetary unit sampling: Improving estimation of the total audit error. *Advances in Accounting*, 25(2), 174-182.
- [9] Hoogduin, L. A., Hall, T. W., & Tsay, J. J. (2010). Modified Sieve Sampling: A Method for Single- and Multi-Stage Probability-Proportional-to-Size Sampling. *AUDITING: A Journal of Practice & Theory*, 29(1), 125-148.
- [10] Hoogduin, L. A., Hall, T. W., Tsay, J. J., & Pierce, B. J. (2015). Does systematic selection lead to unreliable risk assessments in monetary-unit sampling applications? *Auditing: A Journal of Practice & Theory*, 34(4), 85-107.
- [11] Horgan, J. M. (1994). Monetary-unit sampling: An investigation-Vol 1. City University of London.
- [12] Horgan, J. M. (1996). The Moment Bound with Unrestricted Random, Cell and Sieve Sampling of Monetary Units. *Accounting and Business Research*, 26(3), 215-223.
- [13] Horgan, J. M. (1998). A Comparison of Lahiri and Sieve Sampling with Traditional MUS Methods. *Irish Accounting Review*, 5(1), 1-22.
- [14] Institut Akuntan Publik Indonesia. (2021). Standar Audit 530 (Revisi 2021). IAPI.
- [15] Leslie, D. A., Teitlebaum, A. D., & Anderson, R. J. (1979). Dollar-unit sampling: a practical guide for auditors. Copp Clark Pitman.
- [16] Li, X., Brazel, J. F., & Gold, A. (2024). An Unintended Consequence of Full Population Testing on Auditors' Professional Skepticism. Available at SSRN: <https://ssrn.com/abstract=4989842>
- [17] Mulyadi. (2010). Auditing (6th ed., Vol. 1). Salemba Empat.
- [18] Stewart, T. (2012). Technical notes on the AICPA audit guide Audit Sampling. American Institute of Certified Public Accountants, New York, 5-8.
- [19] Stewart, T., & Sunderland, D. (2024). A Bayesian Audit Assurance Model Incorporating Monetary Unit Sampling. Available at SSRN: <https://ssrn.com/abstract=4817734>
- [20] Sukrisno Agoes. (2004). Auditing (Pemeriksaan Akuntan) oleh Kantor Akuntan Publik (3rd ed.). Lembaga Penerbit FE UI.
- [21] Whittington, R., Pany, K., Whittington, O. R., & Pany, K. (2010). Principles of auditing and other assurance services. McGraw-Hill.
- [22] Woodhead, A. D. (1997). The other audit risk: the impact of false rejection on audit planning. *Managerial Auditing Journal*, 12(1), 4-8.
- [23] Wurst, J., Neter, J., & Godfrey, J. (1989a). Comparison of sieve sampling with random and cell sampling of monetary units. *Journal of the Royal Statistical Society: Series D*, 38(4), 235-249.

AUTHORS

Aan Subarkah – Master's student in Statistics, Faculty of Mathematics and Natural Sciences, Hasanuddin University.

Georgina Maria Tinungki – Professor, Department of Statistics, Faculty of Mathematics and Natural Sciences, Hasanuddin University.

Nurtiti Sunusi – Professor, Department of Statistics, Faculty of Mathematics and Natural Sciences, Hasanuddin University.

Georgina Maria Tinungki – georgina@unhas.ac.id.