

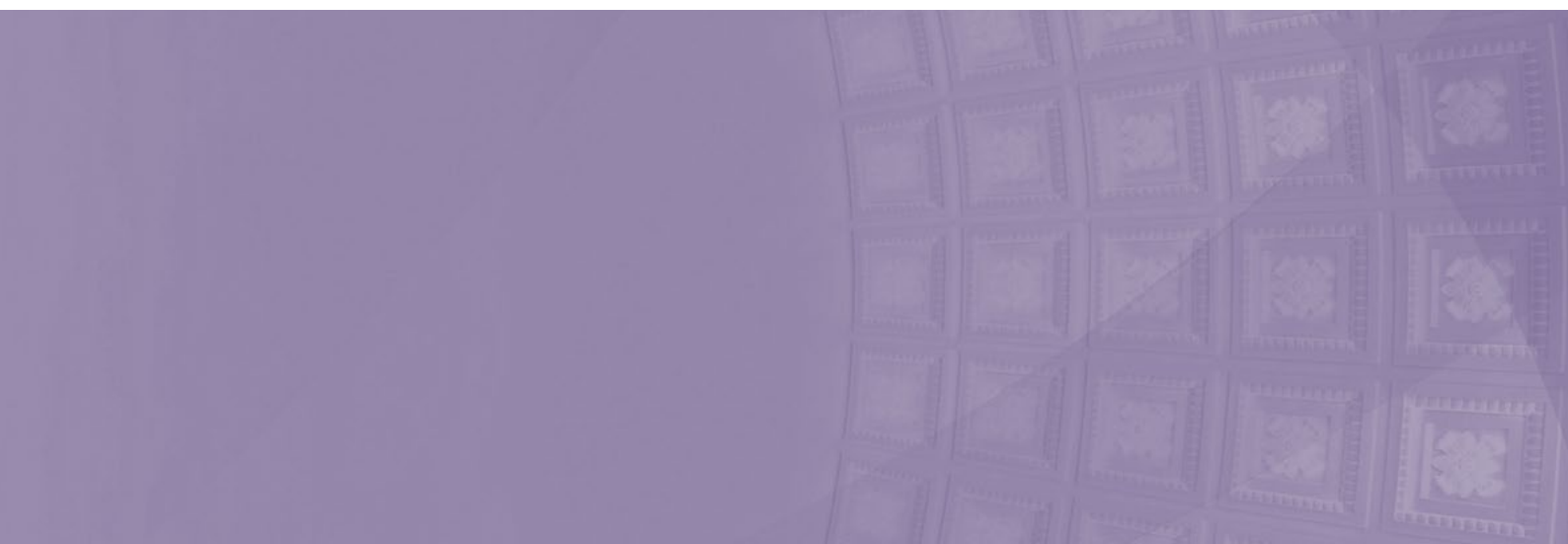


Surveys S-44

Application of Optimised Discretisation Based on the Within-cluster Sum of Squares to Assess the Inherent Risk of Money Laundering and Terrorist Financing

Lucija Brekalo Mandić and Nikolina Maričević

Zagreb, July 2026



The authors would like to thank their families, Vesna Krizmanić and Damir Blažeković for their support.

Author(s)

Nikolina Maričević

Croatian National Bank, Anti-Money Laundering and Terrorist Financing Supervision
Department (Senior Advisor)
nikolina.maricevic@hnb.hr
T. +385 1 45 65 079

Lucija Brekalo Mandić

Croatian National Bank, Anti-Money Laundering and Terrorist Financing Supervision
Department (Chief Associate)
lucija.brekalo-mandic@hnb.hr
T. +385 1 46 90 117

Application of Optimised Discretisation Based on the Within-cluster Sum of Squares to Assess the Inherent Risk of Money Laundering and Terrorist Financing

Abstract

In the context of increasingly demanding regulatory frameworks and the dynamic nature of the financial market, an accurate and objective assessment of inherent money laundering and terrorist financing (ML/TF) risks has become a regulatory imperative. The paper proposes a methodological approach based on the permutation method and the method of minimising within-cluster variance, adapted for the classification of quantitative ML/TF risk indicators into categories using an interval scoring system. The model allows for the empirical evaluation of inherent risk categories – such as customer risk, product and service risk, geographical risk, and distribution channel risk – based on quantitative data provided by supervised entities. A key mathematical challenge of this methodology is the combinatorial complexity arising from the distribution of data across multiple intervals, which requires algorithmic processing and cannot be conducted manually in a reliable way. By applying a permutation-based optimisation approach in combination with least-squares minimisation, the model ensures statistically consistent scoring with significantly reduced subjectivity. Although the method applied in this analysis is not a classical least-squares regression, it is based on the same mathematical principle of quadratic optimisation. Specifically, the method minimises the sum of squared deviations of observations from their group means, with the objective of achieving minimum within-cluster variance. The optimisation is performed over discrete group boundaries rather than over a continuous function, which makes the approach suitable for risk classification and score discretisation rather than regression modelling. The method applies a least-squares-type-quadratic minimisation principle to discretise risk indicators into homogeneous groups by minimising within-cluster variance. This methodology is adaptable for implementation in dynamic Excel environments or Python-based systems, which makes it a practical and scalable solution for supervisory authorities. The model aligns with the evolving requirements of risk-based supervision frameworks.

Keywords: inherent risk, risk assessment, data quality, prevention of money laundering and terrorist financing, permutation method, least-squares method, method of minimising within-cluster variance, algorithmic scoring

JEL codes: C15, C43, C58, C92

Contents

Contents	4
1 Introduction.....	5
1.1 Broader context of risk assessment in supervision.....	5
2 Literature overview.....	7
3 Data.....	10
4 Methodology.....	15
4.1 Inherent risk assessment process.....	18
5 Empirical assessment of inherent risk	20
5.1 Role of permutations and combinatorial complexity	25
6 Discussion.....	27
7 Role of supervisors and supervisory corrections	28
8 Conclusion	29
References	31

1 Introduction

An assessment of inherent risks in anti-money laundering and terrorist financing is an essential segment in the process of supervising financial institutions. Bearing in mind the increasingly complex regulations and dynamic nature of the financial market, it is necessary to adapt the existing risk assessment methods to new requirements that ensure objectivity and precision in analysis. In this context, the use of mathematical models and algorithmic thinking allows for standardised, quantitative and reproducible risk evaluation, significantly reducing space for subjective interpretations and ensuring a more consistent and transparent assessment in line with supervisory objectives. The paper proposes a deterministic method of optimised data distribution into a predefined number of ordinal ratings, i.e. intervals, based on the global within-cluster sum of squares. Although similar to traditional clustering methods such as the k-means algorithm, the proposed approach differs in that it does not use iterative heuristics or initial centroids but considers the entire set of admissible partitions subject to regulatory constraints. The following text shows an empirical example of the assessment of inherent ML/TF risk using the method of minimising within-cluster variance with a predefined number of risk intervals. By applying the method to simulated data, it is possible to understand the risk indicator analysis process and to form a final inherent risk score. The results show the detailed process of assessing inherent risk for 16 supervised entities, including the process of assessing individual risk indicators, assessing risk categories and formulating final scores, as well as the subsequent allocation of entities across risk profiles. This demonstrates its practical application and a methodological framework that allows for precise and objective identification of higher risk entities and supports informed supervisory decision-making.

1.1 Broader context of risk assessment in supervision

Risk assessment is a supervisory process used to identify, analyse and evaluate risks in financial institutions and sectors to assess the likelihood of their materialisation, which ensures that supervisory resources are channelled to areas of higher risk. In essence, risk assessment enables supervisors to make informed decisions. In the context of ML/TF prevention, a risk-based approach to supervision is a key regulatory principle requiring the systematic identification and assessment of risks prior to the application of supervisory measures. This approach involves adjusting supervisory activities to assessed risks. The aim is to enable supervisors to allocate limited resources in an

optimal way to manage ML/TF risks effectively and to ensure that supervisory activities and attention are focused on higher-risk areas. On the other hand, it also ensures that supervisory activities do not impose undue burdens on low-risk sectors and entities, which is key to preserving financial inclusion. Ultimately, this can reduce overall ML/TF risks through a higher level of transparency (FATF). Risk assessment, therefore, serves not only to rank institutions but also to decide on the intensity, depth and focus of supervisory activities. To apply risk-based supervision, supervisors must first understand the exposure of supervised sectors and individual institutions to ML/TF risks. Understanding ML/TF risks at sectoral level is important to prioritise supervisory activities across sectors, while risk assessment of individual supervised entities is used to identify the level of risk of an individual entity within the sector and to steer the level and focus of supervisory engagement within a sector (FATF). An individual risk assessment of the entities subject to money laundering and terrorist financing supervision (hereinafter: individual risk assessment) plays a key role in supervision as it enables supervisors to develop targeted supervisory strategies for each supervised entity, which is a key element in the precautionary approach to ML/TF risk management. Unlike sectoral or national risk assessments, which are based on general levels of threats and vulnerabilities within individual sectors or within a country and provide a broader picture, an individual risk assessment takes into account very specific characteristics and data pertaining to individual supervised entities. This enables the identification of even latent risks that are perhaps not identified at sectoral level but are significant for the specific supervised entities. An individual risk assessment is therefore essential for detailed analysis of the risks that may affect an individual supervised entity. Individual risk assessment methodologies cover a wide range of supervisory approaches, which may vary depending on various factors, such as the country's supervisory framework, the number of supervised sectors and the number of individual supervised entities within each sector. In this context, it is important to understand that a risk assessment relates not only to existing threats but also to the identification of emerging risks. This makes it essential for supervisors to have in place and develop flexible methodologies that allow risk assessments to be continuously adapted and updated to new circumstances. In practice, an individual risk assessment includes an assessment of inherent ML/TF risk and a quality assessment of the control mechanisms implemented by the supervised entity to manage and limit its exposure to ML/TF risks. Inherent risk means the ML/TF risk to which a supervised entity is exposed due to the products and services it offers, its clients, the jurisdiction in which it operates and the distribution channels it uses in its operations, before any mitigating measures or control mechanisms are applied. **Residual risk** is the ML/TF risk to which a supervised entity

remains exposed even after having implemented control mechanisms or mitigating measures (EBA). As illustrated above, an inherent risk assessment is an initial and methodologically crucial phase of a comprehensive risk analysis. Inherent risk describes the underlying risk exposure before control mechanisms and mitigating measures are taken into account, while residual risk is conceptually based on the assessment results. Due to its position at the very beginning of the analytical process, inherent risk assessment is of paramount importance for supervisors, as errors or biases occurring at this stage can be multiplied at later stages of supervision, leading to incorrect prioritisation of supervised entities and ineffective allocation of supervisory resources (Power, 2009). This paper focuses on individual assessment of inherent risk by exploring the proposed methodological approach to inherent risk scoring. Different supervisory authorities apply different methodological approaches to quantifying inherent risk. While some supervisors use quantitative models, others rely on qualitative analysis or combined approaches. However, there is a growing trend towards the use of mathematical methods to allow for a more precise, objective and systematic risk analysis.

2 Literature overview

The requirements for risk-based supervision are defined in the Recommendations of the Financial Action Task Force (the international body which sets global standards and monitors the implementation of measures to combat money laundering, terrorist financing and proliferation financing; hereinafter: FATF).

In its Recommendations¹ and **Risk-Based Supervision** Guidelines, the FATF states that such an approach to supervision encompasses the identification, assessment and understanding of ML/TF risks at sectoral- and supervised entity-level, as well as the application of supervisory measures commensurate with the identified level of risk. Recommendation 1 (R.1) and its interpretative note (INR.1) explain the risk-based approach and Recommendation 2 (R.2) highlights the importance of national co-ordination, including with and among AML/CFT supervisors. R.1 and INR.1 require

¹ FATF Recommendations set out a comprehensive and coherent framework of measures that countries should implement to combat money laundering, terrorist financing, as well as the financing of the proliferation of weapons of mass destruction.

jurisdictions to identify, assess and understand the ML/TF risks and to apply a risk-based approach to mitigate the risks accordingly. This also applies to supervisory activities. The risk-based approach (RBA) set out in R.1 is a foundation for allocating resources and implementing measures to combat ML/TF. The RBA applies in relation to how supervised entities should comply with the regulatory requirements, and how those entities should be supervised. Recommendation 26 (R.26) requires risk-based supervision of financial institutions, while its interpretative note INR.26 recommends that supervisors allocate their supervisory resources according to risk. This requires supervisors to understand the ML/TF risks in their jurisdiction, sector and entities, and to have access to all information relevant to those risks. In the Guidelines on the characteristics of a risk-based approach to anti-money laundering and terrorist financing supervision, and the steps to be taken when conducting supervision on a risk-sensitive basis under Article 48(10) of Directive (EU) 2015/849 (amending the Joint Guidelines ESAs/2016/72), the European Banking Authority (hereinafter: EBA) offers a joint methodological framework for an EU-level risk assessment, without departing from the FATF standards. In the Guidelines, the EBA encourages the application of structured risk assessment methods and integration of data analytics into supervisory practices. As regards the European regulatory framework, Regulation (EU) 2024/1624 of the European Parliament and of the Council of 31 May 2024 on the prevention of the use of the financial system for the purposes of money laundering or terrorist financing (EU) No 1093/2010, (EU) No 1094/2010 and (EU) No 1095/2010 provides for the establishment of an Authority for anti-money laundering and countering the financing of terrorism (AMLA), a central agency responsible for the supervision, coordination and harmonisation of supervision of the systems for preventing money laundering and terrorist financing in the European Union. Pursuant to Article 40, paragraph (2) of Directive (EU) 2024/1640 of the European Parliament and of the Council of 31 May 2024 on the mechanisms to be put in place by the Member States for the prevention of the use of the financial system for the purposes of money laundering or terrorist financing, amending Directive (EU) 2019/1937, and amending and repealing Directive (EU) 2015/849, the AMLA will develop draft regulatory technical standards setting out the benchmarks and a methodology for assessing and classifying the inherent and residual risk profile of obliged entities. While the Regulation does not prescribe specific mathematical methods, it focuses on objective, consistent and data-based risk assessment, implying the need for quantitative approaches to risk assessment. Its provisions also entail the harmonisation of the risk assessment methodology at EU level. This regulatory context opens up scope for the use of statistical methods that allow for a consistent classification and evaluation of inherent risk. Academic literature

very clearly shows that statistical methods can be used in the assessment, classification and discretisation of risks. One of the most prominent examples is the Basel AML Index, which ranks countries according to the risk of money laundering and related financial criminal activities using a combination of 17 indicators from different domains, including the quality of the AML/CFT framework, corruption and fraud risks and financial transparency and standards. For the classification of results into risk categories, this index uses the Jenks natural breaks classification method, which identifies natural differences in the dataset by minimising variance within clusters while maximising differences among them, resulting in categorical risk levels adjusted to actual data distribution rather than arbitrarily defined thresholds. This approach illustrates how statistical methods are applied in practice for the structured determination of risk boundaries in real datasets, which is directly linked to the principles of automatic discretisation of numerical risk indicators. In addition, academic literature increasingly recognises that ML risk assessment cannot be based solely on qualitative or expert-heuristic assessments but requires a certain degree of mathematical formalisation. Krastiņš (2019) assumes that the ML/TF risk is not a classic financial risk with a clearly defined probability structure, but rather a complex concept comprising multiple heterogeneous risk factors. In order for such factors to be consistently aggregated into a single risk score, the author proposes the use of a fuzzy logic approach, where individual risk factors are translated into numerical values and aggregated using mathematical functions. The result of this exercise is a discrete classification across several risk ratings (low, medium, high and non-acceptable). Vagaská et al. also develop a non-linear mathematical and statistical model that describes the behaviour of the Basel AML Index through a set of macroeconomic and socio-economic indicators and employ optimisation algorithms to find combinations of factors where the risk of legalising income from criminal activities is minimal. This approach combines regression methods, non-linear modelling and numerical optimisation to describe the complex risk system at the level of EU member states. In the paper “The anti-money laundering risk assessment: A probabilistic approach” (Journal of Business Research, 2023), Ogbeide et al. show that, in practice, risk assessment often depends on heuristic and procedural approaches, with experts and beginners showing strong cognitive biases and excessive confidence in the assessment of risk distributions and preferring false positives when analysing potentially suspicious transactions of supervised entities. Such findings highlight the limitations of subjective judgement in risk assessment and support the need for formal, structured and data-supported risk assessment methods that reduce the impact of human biases and enhance evaluation consistency. A similar conceptual framework has been applied for decades in

the area of credit risk, where banks and supervisors used structured scoring and classification systems to quantify the inherent credit risk of borrowers. Under Basel II and Basel III regulations, and in particular through the Internal Ratings-Based (IRB) approach, banks have been developing internal models for estimating the probability of default (PD), continuous risk indicators being translated into discrete rating grades (BCBS, 2006; BCBS, 2017). These grades are defined to cover relatively homogeneous groups of debtors, where regulators require that intra-group differences be minimal, while inter-group differences must be clearly expressed and stable over time (BCBS, 2005; Engelmann & Rauhmeier, 2011). Central banks and supervisory authorities further emphasise the importance of transparent and reproducible risk clustering methods, as discrete risk categories are the basis for supervisory decisions, stress tests and macroprudential analyses (ECB, 2018; EBA, 2020). While different statistical techniques are used in practice, their common denominator is the implicit or explicit aspiration to minimise within-group variance and maximise inter-group variance, ensuring that each risk category represents a meaningful and interpretable set of entities with a similar risk profile. The method proposed in this paper builds conceptually on the regulatory and academic practices described but adapts them to the specificities of the supervisory context. The optimisation of within-cluster variance at the level of individual risk indicators and supervised entities allows for a transparent, reproducible and data-based inherent risk scoring. Such an approach strengthens the robustness of assessment at an early stage of the process and reduces the likelihood of initial errors being passed on and amplified at later stages of residual risk analysis. Due to its methodological similarity with approaches already used in credit risk and macroprudential supervision, the proposed methodology may be relevant and applicable for other regulatory and supervisory authorities facing comparable challenges in assessing inherent risk.

3 Data

In accordance with the Anti-Money Laundering and Terrorist Financing Act (Official Gazette 108/2017, 39/2019 and 151/2022), the Croatian National Bank (CNB) conducts supervision over a part of obliged entities in the area of the prevention of money laundering and terrorist financing. When planning and conducting supervision, the CNB is required to apply an approach based on the assessment of money laundering and terrorist financing risks, in order to ensure the proportionality of supervision. In this

respect, the exercise of supervision necessarily entails a systematic individual risk assessment of supervised entities, based on an inherent risk assessment. In supervisory practice, inherent risk assessment relies on information gathered through regular questionnaires, supervisory reports, inspection findings and other relevant sources. Based on the collected data, risk indicators within individual inherent risk categories are created. To illustrate the proposed methodology for inherent risk assessment, the empirical part of this paper is based on simulated quantitative data. The simulated dataset comprises a hypothetical sample of supervised entities and provides a simplified presentation of quantitative risk indicators. The selected approach enables a presentation of methodological steps, comparison with alternative clustering techniques, and discussion of the benefits and limitations of the proposed method. According to the ML/TF Risk Factors Guidelines (EBA/GL/2021/02), supervisors should consider risk factors related to the characteristics of customers, geographical areas, the types of products and services, and delivery channels. In line with the risk-based approach, the supervisory framework takes into account different sources of risk, including the characteristics of products and services, customer structure, the geographical reach of operations and the distribution channels used by the institution. The assessment of inherent risk is most often based on categories such as:

- customer profile (e.g. politically exposed persons, residence);
- type of products and services (e.g. private banking, virtual assets);
- geographical risk (e.g. operations with high-risk countries);
- distribution channels (e.g. use of intermediaries).

Inherent risk assessment does not seek to determine a universal or absolute risk value, but rather focuses on establishing a flexible, quantitatively based framework that allows an understanding of the risk exposure. Such a framework is based on a comparative analysis of the supervised entities.

As mentioned earlier, the inherent risk of money laundering and terrorist financing relates to the original, uncontrolled risk arising from the nature of the supervised entity's business and the characteristics of its customers, its products and the markets in which it operates and the distribution channels it uses, without taking into account mitigating measures; therefore, such an assessment is a mandatory component of residual risk assessment.

In this context, an inherent risk assessment is based on a quantitative assessment of risk indicators from four categories, including customer risk, product and service risk,

geographical exposure and distribution channels, using the method of minimising within-cluster variance. The objective is to assign a score of 1 to 4 to each quantitative data point belonging to the risk indicator being assessed, according to the following categorisation:

Score 1	Low risk
Score 2	Medium risk
Score 3	Increased risk
Score 4	High risk

EBA **regulatory guidelines** suggest the categorisation of the overall inherent risk level into four categories: less-significant risk (1), moderately significant risk (2), significant risk (3) and very significant risk (4). This approach directly supports the implementation of quantitative and mathematically based methods, such as the method of minimising within-cluster variance, which enable an objective and transparent evaluation of inherent risk indicators. However, those guidelines do not require the use of specific statistical or mathematical algorithms to assess the inherent risk of money laundering and terrorist financing, but impose principle-based requirements for a standardised, documented and proportionate method within a risk-based approach. Therefore, the method proposed in this paper is not a regulatory technique, but rather a methodological proposal that aligns with the above guidelines.

For the purpose of illustrating the methodology, a hypothetical sample of 16 units of analysis is considered, with each unit representing one entity. Quantitative risk indicators were generated for each unit, corresponding to individual inherent risk categories. In general, the indicators capture and reflect the characteristics related to customer risk, product and service risk, geographical exposure and distribution channels. Table 1 shows data for risk indicators used to estimate customer risk, Table 2 shows data for risk indicators used to estimate product /service risk, while Tables 3 and 4 show data for risk indicators used to assess geographical risk and distribution channels risk, respectively. In each of the tables, $S_1, \dots, S_{16}$ indicate the supervised entities and $I_1, \dots, I_7$ stand for risk indicators.

Table 1 Overview of risk indicators in the customer risk category

	I_1	I_2	I_3	I_4	I_5	I_6	I_7
S_1	1	0.18	0.23	21137.00	1	89	2113

S_2	2	0.18	0.45	39479.00	2	112	24211
S_3	2	0.42	0.55	17682.00	2	220	9479
S_4	3	0.51	0.13	16856.00	2	20	16845
S_5	4	0.23	0.75	56301.00	3	133	38943
S_6	30	0.53	0.32	2389.00	3	161	31577
S_7	33	0.34	0.82	49605.00	4	150	46309
S_8	50	0.37	0.12	41648.00	4	170	53675
S_9	57	0.66	0.58	6543.00	10	171	75773
S_{10}	75	0.85	0.43	126010.00	11	180	61041
S_{11}	87	0.46	0.50	22498.00	11	181	68407
S_{12}	110	0.19	0.13	45513.00	12	200	83139
S_{13}	600	0.91	0.32	809438.00	20	191	97871
S_{14}	620	0.28	0.44	1010331.00	21	192	90505
S_{15}	700	0.10	0.93	593177.00	22	182	105237
S_{16}	750	0.11	0.48	1377152.00	23	190	112603

Source: Authors' calculations.

Table 2 Overview of risk indicators in the risk category of products and services

	I_1	I_2	I_3	I_4	I_5	I_6
S_1	0.10	0.05	0.18	1130	9478	11845
S_2	0.20	0.25	0.15	23211	31877	38643
S_3	0.22	0.18	0.30	46309	56675	62041
S_4	0.35	0.40	0.28	68407	70173	84139
S_5	0.50	0.45	0.52	90505	97871	104237
S_6	0.42	0.55	0.50	122603	1130	9478
S_7	0.75	0.80	0.70	11845	23211	31877
S_8	0.28	0.33	0.37	38643	46309	56675
S_9	0.45	0.50	0.48	62041	68407	70173
S_{10}	0.55	0.60	0.50	84139	90505	97871
S_{11}	0.52	0.45	0.50	104237	122603	1130
S_{12}	0.60	0.55	0.58	9478	11845	23211

S_{13}	0.62	0.58	0.55	31877	38643	46309
S_{14}	0.65	0.60	0.62	56675	62041	68407
S_{15}	0.80	0.85	0.75	70173	84139	90505
S_{16}	0.78	0.82	0.80	97871	104237	122603

Source: Authors' calculations.

Table 3 Overview of risk indicators in the geographical risk category

	I_1	I_2	I_3	I_4	I_5
S_1	1700835.98	120000.46	0.47	1708345.22	0.01
S_2	1380969.21	13189967.45	0.41	3621971.65	0.18
S_3	2020702.75	26259934.44	0.50	5535598.09	0.13
S_4	421368.9	39329901.44	0.50	22031.83	0.20
S_5	741235.67	52399868.43	0.00	432594.26	0.25
S_6	1321345.52	65469835.43	0.17	1026406.08	0.22
S_7	1061102.44	78539802.42	0.53	1070469.74	0.32
S_8	3620036.60	91609769.41	0.20	2346220.7	0.29
S_9	2121123.78	104679736.41	0.23	2118907.65	0.36
S_{10}	2340569.52	117749703.40	0.32	2984096.17	0.39
S_{11}	2660436.29	130819670.39	0.59	3167345.56	0.43
S_{12}	2980303.06	143889637.39	0.56	4897722.61	0.46
S_{13}	3300169.83	156959604.38	0.62	4259847.13	0.50
S_{14}	4259770.14	170029571.38	0.68	6173473.56	0.53
S_{15}	3939903.37	183099538.37	0.65	1070469.74	0.57
S_{16}	4579636.91	196169505.36	0.65	6811349.04	0.60

Source: Authors' calculations.

Table 4 Overview of risk indicators in the distribution channels category

	I_1	I_2	I_3	I_4	I_5
S_1	22	0.21	12389	0.113	1130
S_2	27	0.24	16543	0.112	9478

S_3	100	0.27	136010	0.213	11845
S_4	111	0.25	16856	0.223	23211
S_5	322	0.30	18682	0.312	31877
S_6	10	0.29	21436	0.322	46309
S_7	15	0.32	22598	0.413	38643
S_8	31	0.34	49479	0.424	56675
S_9	33	0.57	61648	0.752	68407
S_{10}	49	0.62	65513	0.451	62041
S_{11}	41	0.70	69605	0.486	70173
S_{12}	122	0.67	76301	0.511	84139
S_{13}	113	0.76	909438	0.551	90505
S_{14}	201	0.66	593177	0.582	97871
S_{15}	210	0.77	1020331	0.782	104237
S_{16}	423	0.80	1377152	0.930	122603

Source: Authors' calculations.

4 Methodology

Assessing inherent ML/TF risk requires a structured and quantitatively-based approach, to allow for a consistent and reproducible evaluation of the various risks across supervised entities. Given the complexity and the significant number of inputs, which include diversity of customers, products, geographical areas and distribution channels, it is appropriate to apply a method that reduces the impact of subjective interpretations and allows objective risk classification. This paper describes the method of minimising within-cluster variance under certain conditions in order to convert quantitative data for each individual risk indicator within a specific risk category into discrete risk scores on a scale of 1 to 4. In this context, risk discretisation refers to the process of transforming quantitative risk indicators into a final number of discrete, ordinal structured risk ratings, with the aim of preserving the information content of the data and ensuring consistent interpretation of the results in the supervisory and regulatory context. The objective is to find optimal data distribution in four intervals so as to minimise the total sum of squared deviations of each data point from the arithmetic mean of the

corresponding interval. Such minimisation ensures the minimum within-cluster variance and therefore constitutes a mathematically justified way of forming homogeneous groups that reflect different levels of risk. At the same time, this method is based on empirical indicators (as it assesses quantitative data), mathematical and statistical validation (as the scoring is based on the application of permutations and variance minimisation) and algorithmic structure (as the method is programmable and subject to automated processing, which makes it suitable for implementation within contemporary supervisory frameworks based on data and algorithmic decision-making). The methodology presented in this paper is based on the deterministic algorithmic process of grouping quantitative data into a predefined number of intervals, whereby intervals are interpreted as the corresponding scores, i.e. inherent risk levels. The goal is to ensure an objective, reproducible and mathematically sound allocation to risk ratings. For each risk indicator, a set of quantitative observations is considered:

$$\{x_1, x_2, \dots, x_n\},$$

where n is the number of supervised entities. Data are first sorted in descending order and marked as $x_{(1)} \geq x_{(2)} \geq \dots \geq x_{(n)}$. Then considered are all possible set partitions into a predefined number of intervals (in this paper, four intervals), provided that each interval contains at least one observed value. For each possible partition, the within-interval sum of squares (WCSS) is calculated, defined as:

$$WCSS = \sum_{j=1}^k \sum_{x_i \in I_j} (x_i - \bar{x}_j)^2,$$

where I_j denotes the j -th interval and $\bar{x}_j$ its arithmetic mean. The optimum partition is the one that minimises the total value of the WCSS. It is necessary to examine all possible partitions of the data in four intervals, which means that all possible partitions of the sorted set in four intervals need to be tested for all quantitative data, followed by the application of the method of minimising within-cluster variance in order to determine the appropriate data partition and the score for each indicator and supervised entity. First of all, when assessing an individual inherent indicator, appropriate quantitative data relating to the ML/TF system for supervised entities are taken into account and data are sorted from the smallest to the largest, after which the conditions of a technical and logical nature that ensure the proper operation of the algorithm are set. The conditions are determined automatically as follows: k is the number of intervals/levels and n is the number of supervised entities reporting quantitative data for

a specific inherent risk indicator or category. All data are sorted from the smallest to the largest, i.e. they are placed on the number line (real x -axis) and the following conditions are determined:

$$t_1, t_2, \dots, t_{k-1} \in \{1, 2, 3, \dots, n-1\}; t_1 < t_2 < \dots < t_{k-1},$$

which set the boundaries between the intervals/groups so that t_j represents the number of data points to the left of the boundary between the j -th and the $j+1$ interval.

Furthermore, once all possible partitions of the data have been defined within k intervals, the global minimum of the function is determined in accordance with the conditions set out above:

$$L(t_1, \dots, t_{k-1}) = \sum_{j=1}^k \sum_{x_i \in I_j} (x_i - \bar{x}_j)^2,$$

where $I_j = \{x_{(t_{j-1}+1)}, \dots, x_{(t_j)}\}$ is the set of all points belonging to the j -th interval and $\bar{x}_j$ is their average. Unlike iterative cluster methods, this approach explicitly searches the space of possible partitions and provides the global minimum of the target function for the specified number of intervals. This removes the need for initialisation, heuristics or additional decision-making rules during the clustering process. The objective is to find the optimal division of a set of observations into a predefined number of intervals k (in this paper $k = 4$), so as to minimise the total sum of squared deviations of observations from the arithmetic means of the respective intervals. This ensures the minimum within-cluster variance and the formation of homogeneous groups representing different levels of inherent risk. The distribution of the dataset in four intervals is defined by the three boundary indices $1 \leq t_1 < t_2 < t_3 < n$, which determine the boundaries between adjacent intervals. The following intervals are defined:

$$I_1 = \{x_{(1)}, x_{(2)}, \dots, x_{(t_1)}\},$$

$$I_2 = \{x_{(t_1+1)}, \dots, x_{(t_2)}\},$$

$$I_3 = \{x_{(t_2+1)}, \dots, x_{(t_3)}\},$$

$$I_4 = \{x_{(t_3+1)}, \dots, x_n\}.$$

For each admissible combination of indices (t_1, t_2, t_3) , the value of the target function, defined as the total within-cluster sum of squared deviations, is calculated:

$L(t_1, t_2, t_3) = \sum_{j=1}^4 \sum_{x_i \in I_j} (x_i - \bar{x}_j)^2$, where $\bar{x}_j$ denotes the arithmetic mean of observations in interval I_j . The algorithm searches all admissible combinations of indices (t_1, t_2, t_3) and selects the one that minimises the value of the target function. Once the optimal partition has been determined, the boundaries between adjacent intervals are defined as arithmetic means of the endpoint observations of adjacent intervals. The first interval has no lower limit and the last interval has no upper limit. The inherent risk scoring scale in the range from 1 to 4 is defined as follows:

Score 1 = Low risk	$\langle -\infty, \frac{x_{(t_1)} + x_{(t_1+1)}}{2}]$
Score 2 = Medium risk	$\left[\frac{x_{(t_1)} + x_{(t_1+1)}}{2}, \frac{x_{(t_2)} + x_{(t_2+1)}}{2} \right)$
Score 3 = Increased risk	$\left[\frac{x_{(t_2)} + x_{(t_2+1)}}{2}, \frac{x_{(t_3)} + x_{(t_3+1)}}{2} \right)$
Score 4 = High risk	$\left[\frac{x_{(t_3)} + x_{(t_3+1)}}{2}, \infty \right)$

4.1 Inherent risk assessment process

In the proposed model for assessing inherent ML/TF risk, the assessment is performed through a strictly defined sequence of analytical and algorithmic steps. The methodological framework is based on the quantitative processing of data collected directly from supervised entities. These data are used as an input variable for predefined ML/TF risk indicators, ensuring their relevance and measurability. Each indicator represents a measurable response to the level of potential risk exposure and is associated with one of the four inherent risk categories: customer risk, product and service risk, geographical exposure risk and distribution channel risk. Within each inherent risk category, several indicators make up a data set that is individually evaluated using the method of minimising within-cluster variance and then aggregated. This allows for a detailed and differentiated assessment of each risk category. The most important segment of inherent risk assessment is the creation of risk indicators, which provide the

basis for further quantitative analysis and assessment of the risk exposure level of the supervised entity. Precisely defined indicators enable a standardised, consistent and objective measurement of specific risk aspects. The regulatory framework for the prevention of money laundering and terrorist financing does not prescribe indicators or the number of indicators to be used in a particular category of inherent risk; instead, it provides only general guidelines. In practice, inherent risk categories, such as customer risk, product and service risks, distribution channel risk and geographical risk, are typically operationalised depending on the complexity of the business model of the sector of supervised entities, as well as the size of the sector itself. Quantitative data used to define most risk indicators are collected from supervised entities on an annual basis or more frequently if necessary. The risk indicators defined are then evaluated using the method of minimising within-cluster variance by means of algorithmic processing implemented in the computer environment; such evaluation determines the optimal boundaries between the four levels of risk with each quantitative data point allocated to one of the possible four risk levels, i.e. four intervals, the data being assigned the corresponding scores based on their interval (1 – 4), each of which results in a numerical risk score: 1 – low risk, 2 – medium risk, 3 – increased risk, 4 – high risk. This means that an assessment and calculation of the corresponding scores have been made for each supervised entity covered by the risk assessment and for each risk indicator. After calculating risk indicator scores, the final score for each risk category is calculated. The scores of intra-category indicators are aggregated by an arithmetic mean, after which the aggregated scores are translated into discrete risk scores from 1 to 4, again by applying an algorithm based on the method of minimising within-cluster variance, which results in a single score of the inherent risk category. In practice and in the literature, there are alternative approaches to the aggregation of indicators, such as weighted average, maximum value, quantile and non-linear aggregation, as regulatory guidelines do not prescribe a single function but allow supervisor discretion. The single scores of risk categories are then weighted in accordance with predefined weight values. Each of the four risk categories bears a certain weight in the overall risk score, reflecting its relative importance in the entity's total risk. In the process of aggregating individual inherent risk categories to a single overall score, weights are applied

reflecting the relative importance of individual risk sources, such as customer risk, product and service risk, distribution channel risk and geographical risk. The purpose of weighting is not purely technical in nature, but constitutes a key mechanism for operationalising the risk-based approach, which allows the methodology to be adapted to the specificities of the national financial system and dominant ML/TF typologies. Regulatory guidelines do not prescribe predefined weight values for individual risk indicators and categories but leave room for competent authorities to determine, within their legal powers and based on supervisory judgement, the weights appropriate to the identified risks, the results of the national risk assessment and the structure of supervised entities. Such flexibility allows regulators to ensure the proportionality and effectiveness of supervision, while preserving methodology sensitivity to changes in the risk environment. For this reason, this paper considers weighting at a conceptual level, while the empirical part is based on illustrative and simulated values, ensuring scientific reproducibility. After the application of weights and aggregation of categorical risk scores, an overall numerical value of inherent risk is obtained, to which a qualitative risk score (e.g. low, medium, increased risk, high risk) is again assigned using an algorithm based on the method of minimising within-cluster variance.

5 Empirical assessment of inherent risk

The empirical application of the proposed methodology was performed on a simulated dataset comprising 16 hypothetical supervised entities. For each supervised entity, selected risk indicators are assessed within four categories of inherent risk, where the scope of the analysed data covers the following components:

- customer risk: 7 indicators;
- product and service risk: 6 indicators;
- geographical risk: 5 indicators;
- distribution channel risk: 5 indicators.

The assessed simulated data presented in the tables in the Data section are designed to reflect the often-observed heterogeneity of data sizes. The method of minimising within-cluster variance was used for each indicator, with indicator values grouped into

four risk ratings as a result of the assessment. This ensures an objective assignment of discrete scores from 1 to 4, with minimisation of intra-group variability and maximisation of inter-group differences. The final score for inherent risk categories, and ultimately the final inherent risk score, is obtained through the calculation described in Section 4.1. The illustration of the algorithmic score calculation by using the method of minimising within-cluster variance for an individual risk indicator will be shown using the example of *Indicator 1* in Table 1. The first step covers data sorting from the smallest to the largest, after which the following conditions are set: $n = 4$ is the number of intervals, i.e. score categories, and $b = 16$ is the number of supervised entities reporting data for *Indicator 1* in the customer risk category. Then determined are the numbers $t_1, t_2, t_3 \in \{1, 2, 3, \dots, 15\}$; $t_1 < t_2 < \dots < t_3$, which define the boundaries between the intervals so that t_i represents the number of data to the left of the boundary between the i -th and the $i + 1$ interval. For example, if $t_1 = 2$, the first two data points belong to interval 1, and the third largest data point belongs to the next interval. The process continues with the generalisation of the principle that within individual intervals the data must be as close as possible, which ensures the best partition of the data being assessed. In the present case, there are 451 possible data partitions in four intervals, provided that no interval is empty. This means that there is a minimum of one and a maximum of 13 data in an interval, provided that the data in each subsequent interval are bigger than the data in the previous interval. The global minimum function $L(t_1, \dots, t_3) = \sum_{i=1}^4 \sum_{q \in x_i} (q - \bar{x}_i)^2$, is then determined for 451 possible data partitions in four intervals, where $x_i = x_i(t_i, t_{i-1})$ is the set of all data belonging to the i -th interval, $\bar{x}_i$ is their average, and q represents the data within a given interval. In this case, the global minimum of the function is 4984.23, determined by the following data partition:

1st interval = (1, 2, 2, 3, 4, 30, 33)

2nd interval = (50, 57, 75, 87, 110)

3rd interval = (600, 620)

4th interval = (700, 750).

The data, i.e. the supervised entities whose data are in the 1st interval, are assigned a score of 1, representing the minimum probability of ML/TF risk, a score of 2 is assigned to the second interval, a score of 3 is assigned to the 3rd interval and a score of 4 is given to the 4th interval containing data indicating the highest likelihood of money

laundering risk. The third row of Table 5 shows the scores for *Indicator 1* in the customer risk category.

Table 5 Overview of the scores for *Indicator 1* for 16 supervised entities

S_1	S_2	S_3	S_4	S_5	S_6	S_7	S_8	S_9	S_{10}	S_{11}	S_{12}	S_{13}	S_{14}	S_{15}	S_{16}
1	2	2	3	4	30	33	50	57	75	87	110	600	620	700	750
1	1	1	1	1	1	1	2	2	2	2	2	3	3	4	4
Low risk							Medium risk					Increased risk		High risk	

Note: The highlighted area indicates the dataset relevant for the assessment of the inherent risk of hypothetical entities.

Source: Authors' calculations.

Scores for all indicators within individual risk categories are determined in the same way. For risk indicators within the customer risk category, the scores shown in Table 6 were obtained using the described scoring method.

Table 6 Overview of risk indicator scores for the customer risk category

	I_1	I_2	I_3	I_4	I_5	I_6	I_7
S_1	1	1	2	1	1	2	1
S_2	1	1	3	1	1	2	1
S_3	1	2	3	1	1	4	1
S_4	1	3	1	1	1	1	1
S_5	1	2	4	1	2	2	2
S_6	1	3	2	1	2	3	2
S_7	1	2	4	1	2	3	2
S_8	2	2	1	1	2	3	2
S_9	2	3	3	1	3	3	2
S_{10}	2	4	3	1	3	3	3
S_{11}	2	3	3	1	3	3	3
S_{12}	2	1	1	1	3	4	3

S_{13}	3	4	2	3	4	4	4
S_{14}	3	2	3	3	4	4	4
S_{15}	4	1	4	2	4	3	4
S_{16}	4	1	3	4	4	4	4

Source: Authors' calculations.

After calculating risk indicator scores, the final score for each risk category is calculated. Using the formula

$$R_{supervised\ entity,i} = \frac{1}{m} \sum_{j=1}^m r_{ij},$$

where r_{ij} is the score of the j -th indicator for the i -th supervised entity, and m is the number of indicators, an aggregated score of the inherent risk category is obtained. The aggregate scores are then converted into discrete risk scores from 1 to 4, again by using the method of minimising within-cluster variance, which results in a single score of the inherent risk category. The aggregated scores and final scores for the 16 supervised entities for the customer risk category resulting from the application of the above formula are shown in Table 7.

Table 7 Overview of the aggregated and final scores for the customer risk category

Overview of the aggregated and final scores for the customer risk category															
S_1	S_4	S_2	S_3	S_8	S_5	S_6	S_7	S_{12}	S_9	S_{11}	S_{10}	S_{15}	S_{14}	S_{13}	S_{16}
1.29	1.29	1.43	1.86	1.86	2.00	2.00	2.14	2.14	2.43	2.57	2.71	3.14	3.29	3.43	3.43
1	1	1	2	2	2	2	2	2	3	3	3	4	4	4	4

Source: Authors' calculations.

The final scores for the remaining risk categories are also determined in the same way. Based on the risk indicators and the corresponding data shown in Table 2, Table 3 and Table 4 of the Data section, and using the process described in this section for the calculation of the customer risk category score, final scores are obtained for the risk categories of products and services, geographical area and distribution channels, as shown in Table 8.

Table 8 Overview of the final scores of inherent risk categories

Overview of the final scores for the customer risk category															
S_1	S_4	S_2	S_3	S_8	S_5	S_6	S_7	S_{12}	S_9	S_{11}	S_{10}	S_{15}	S_{14}	S_{13}	S_{16}
1	1	1	2	2	2	2	2	2	3	3	3	4	4	4	4
Overview of the final scores for the risk category of products and services															
S_1	S_4	S_2	S_3	S_8	S_9	S_6	S_5	S_{10}	S_{11}	S_{12}	S_{13}	S_{14}	S_7	S_{15}	S_{16}
1	2	2	2	2	2	3	3	3	3	3	3	3	4	4	4
Overview of the final scores for the geographical risk category															
S_5	S_4	S_6	S_1	S_2	S_7	S_3	S_8	S_9	S_{10}	S_{11}	S_{12}	S_{15}	S_{13}	S_{14}	S_{16}
1	1	1	2	2	2	2	2	2	2	2	3	3	3	4	4
Overview of the final scores for the distribution channel risk category															
S_1	S_2	S_3	S_4	S_6	S_7	S_8	S_5	S_{10}	S_{11}	S_9	S_{12}	S_{14}	S_{13}	S_{15}	S_{16}
1	1	1	1	2	2	2	3	3	3	3	3	4	4	4	4

Source: Authors' calculations.

Determination of the final scores for individual inherent risk categories is followed by the aggregation of these scores into a single inherent risk score of the supervised entity. In this case, aggregation is made using a weighted sum, where the weights reflect the relative importance of individual risk categories in the overall score. Simulated weights of equal value were used in this paper (25% per category). This approach allows for a transparent presentation of the methodological framework and calculation mechanics and only serves an illustrative and analytical purpose. The resulting weighted sums for each supervised entity are then classified into the final inherent risk ratings using the within-cluster sum of squares algorithm, which ensures a consistent, objective and data-based distribution of entities according to risk profiles. The mathematical representation of the final score for entity i is calculated using the following formula:

$$IR_i = \sum_{k=1}^4 w_k \cdot R_{ik},$$

where R_{ik} is the final score of category k for entity i , and w_k is the weight for risk category k . In this example, it follows that $w_1 = w_2 = w_3 = w_4 = 0.25$. By inserting the final score values for the risk categories shown in Table 9 into the above formula, aggregate results are obtained and are then again classified into four risk intervals using the within-cluster sum of squares algorithm, which provides the final inherent risk scores in a range from 1 to 4. The aggregated results and the corresponding final inherent risk scores shown in Table 8 represent the distribution of supervised entities

across intervals, indicating the level of the supervised entity's exposure to inherent ML/TF risk.

Table 9 Overview of the final inherent risk scores for 16 supervised entities

Overview of the final inherent risk scores															
S_4	S_1	S_2	S_3	S_6	S_8	S_5	S_7	S_9	S_{10}	S_{11}	S_{12}	S_{13}	S_{15}	S_{14}	S_{16}
1.25	1.25	1.5	1.75	2	2	2.25	2.5	2.5	2.75	2.75	3	3.25	3.75	3.75	4
1	1	1	2	2	2	2	3	3	3	3	3	3	4	4	4

Source: Authors' calculations.

The results given in Table 9 show a clear differentiation of supervised entities according to inherent risk intensity. Entities with extremely high quantitative indicator values are consistently classified into higher risk ratings, while entities with lower values achieve stable low scores across all indicators. Such a structure confirms the consistency and interpretability of the proposed methodology and its applicability in the context of supervisory planning and prioritisation. Based on obtained quantitative and qualitative inherent risk scores, supervisors prioritise supervision, determine the intensity and scope of supervisory activities, allocate resources and define the frequency of monitoring and other measures in respect of entities with a higher ML/TF risk profile.

5.1 Role of permutations and combinatorial complexity

Application of the within-cluster sum of squares (WCSS) in the assessment of inherent risk indicators and categories also entails significant combinatorial complexity. To ensure an appropriate distribution of data, first analysed is the combinatorial complexity of the data being assessed. As combinatorial complexity is a concept often used in computing, graph theory, optimisation and statistics, it is applicable in this case because the solution space is discrete but very large and the target function (WCSS) is continuous and differentiable. As such, the problem belongs to the so-called NP class – there is no known polynomial algorithm for this problem without trying all combinations. The permutation method and the least-squares method have traditionally been used for regression analysis and classification. In this context, these mathematical methods are used for the classification of risk categories and indicators, with the aim that all the quantitative data being evaluated are mapped into four ratings that best “adhere” to the data, which is achieved by minimising within-cluster variance that ensures that the data within each interval are as homogeneous as possible. One of the key challenges in an objective and quantitative assessment of inherent ML/TF risk is the

proper classification of a large number of quantitative indicators into a defined number of risk intervals, such as scores 1 to 4, which are here used as an example of ranking risk levels. As this is not an arbitrary classification but rather an optimal division according to a mathematical criterion, we face the problem of combinatorial explosion. If we have n quantitative data for only one risk indicator and we want to divide them into four intervals, provided that intervals are non-empty and contiguous, we face a problem of setting three cut-off points within $n - 1$ possible positions. The number of such combinations can be expressed by a binomial coefficient:

$$\sum_{i=1}^{n-3} \sum_{j=i+1}^{n-2} \sum_{k=j+1}^{n-1} 1 = \binom{n-1}{3}.$$

In the example presented in this paper, where for one risk indicator there are $n = 16$ observations with 15 potential positions to delineate between intervals, and taking into account the fact that three cut-off points are required for the formation of four ratings, the total number of possible partitions is equal to the following number of combinations:

$$\binom{15}{3} = \frac{15 \cdot 14 \cdot 13}{3 \cdot 2 \cdot 1} = 455.$$

After the exclusion of trivial edge combinations leading to blank intervals, 451 relevant partitions are obtained. This combinatorial complexity clearly illustrates that even if one risk indicator is evaluated, the issue of sorting data into a predefined number of intervals is not trivial and requires a systematic optimisation approach. Computer implementation of this methodological approach can be made by a variety of analytical tools, such as Excel or programming languages dedicated to data processing (e.g. Python). A combination of Sort, Average and Offset functions can be used in Excel to generate all possible partitions and automatically calculate within-cluster variance for each partition. In Python, the method can be programmed using standard libraries such as NumPy and Pandas, where quantitative data are sorted and then all possible combinations of limits are generated by combinations of functions from the itertools module. After calculating the within-cluster variance for each partition, a partition with a minimum WCSS value is automatically selected. Such implementation enables automated and reproducible application of the method to a large number of indicators and supervised entities, minimising subjectivity and the need for manual interventions.

6 Discussion

The proposed approach uses the same target function as the k-means algorithm (within-cluster sum of squares), but unlike the standard k-means clustering, this paper addresses the problem of optimal distribution of one-dimensional quantitative data into a given number of groups with the limitation that groups must be consecutive intervals, which leads to differences in algorithmic approach and interpretation of results. While the standard k-means algorithm uses an iterative heuristic process in which the target function converges into a local minimum and depends on the initialisation of the cluster centres, this paper addresses the distribution of one-dimensional quantitative indicators in an exact manner, by searching all admissible partitions, which always leads to the global minimum and the corresponding optimal data distribution for a predefined number of ratings. Such a structure enables a simple mapping of quantitative indicators into discrete risk scores. In the context of supervisory assessments, the number of risk ratings is often normatively predefined (e.g. four ratings), eliminating the need for heuristic selection of the number of clusters. Application of the standard k-means method minimises the WCSS over a set of all possible observation partitions in k clusters, without additional structural constraints:

$$\min_{\{C_1, \dots, C_k\}} \sum_{j=1}^k \sum_{x \in C_j} \|x_i - \mu_j\|^2,$$

where k indicates the number of intervals, C_j a set of observations and μ_j centroids that are iteratively assessed, with the optimal solution depending on the initialisation of the cluster centres. By applying the method of minimising within-cluster variance with contextual conditions and a predefined number of risk intervals described in this paper, interval boundaries were determined based on the data structure and minimisation of dispersion within the interval, with the cut-off points being set where the differences between the data are actually significant. The obtained intervals clearly separate low, medium and extremely high indicator values. While the methods share the objective of minimising within-cluster-variance, the proposed approach can provide greater interpretability in the context of supervisory assessment of inherent risk, especially in the presence of asymmetric distributions and extreme values.

In addition, the proposed methodology for inherent risk assessment is relative by nature, meaning that the allocation to risk categories is based on the distribution of observed values within the analysed sample rather than on predefined absolute risk thresholds. As this is the general characteristic of discretisation and clustering methodologies, it is

important to emphasise that the objective of the proposed methodology is not to determine the absolute level of ML/TF risk on its own, but to provide a consistent, objective and reproducible framework for the relative differentiation and prioritisation in channelling supervisory resources within a particular group of supervised entities and the time horizon. In situations where the distribution of risk indicators is very narrow or close to being uniform, the algorithm still generates a predefined number of categories, but the actual differences between entities are small. Such cases can be identified in practice by monitoring dispersion measures (e.g. variance, standard deviation or interquartile range) before and after the clustering process. Low dispersion points to a homogeneous population and may signal the need for a more cautious interpretation of the results, a possible reduction in the number of intervals or complementation of the relative classification by absolutely defined risk thresholds and, finally, the implementation of supervisory judgement including qualitative information and contextual indicators. The contribution of the methodology as such lies in the possibility of developing a deterministic process for ranking and discrete classification of one-dimensional, asymmetric quantitative datasets, with a predefined number of intervals and preserving the ordinal meaning of the results. Under such a hybrid approach, the relative nature of the algorithm does not lead to automatic overestimation or underestimation of risks, but rather provides a transparent, standardised and reproducible analytical basis for channelling supervisory resources.

7 Role of supervisors and supervisory corrections

Although the methodology based on minimising within-cluster variance enables an objective, algorithmically consistent and quantitatively precise risk classification, a final score of inherent risk may include supervisory competence and discretionary judgement, where there are specific and documented reasons for doing so, such as regulatory warnings, activities not covered by the algorithm (e.g. the emergence of new risks) or doubts about data quality and accuracy. Supervisors play a key role in ensuring that risk assessment models are not only technically sound but also adapted to the specificities of supervised entities and emerging ML/TF risks. Supervisory intervention is therefore necessary not only for the validation and correction of results, but also for the assessment of input data quality and methodology efficiency in practice. This ensures that technology does not replace professional supervision but facilitates it, thereby increasing the system's resilience to complex and evolving risks, in line with

international best practices. In the context of automated systems for assessing inherent ML/TF risks, it is particularly relevant to stress the importance of input data quality. More specifically, incorrect or incomplete data may lead to the assignment of inadequate scores, which can ultimately lead to erroneous conclusions about the level of risk. Therefore, it is essential that supervisors implement internal controls to ensure data accuracy and completeness, thereby enabling the proper functioning of automated risk assessment systems. For example, if data that formally meet the required format have been submitted for a particular indicator, but the supervision shows that they are based on inadequate classifications or incorrect records, supervisors may correct the score. Also, in situations where an algorithm is unable to identify contextual risks (e.g. the emergence of new risks or product forms), expert intervention provides a necessary correction. However, such corrections cannot be arbitrary. They must be based on supervisory findings, qualitative elements not fully quantified in the model, or new insights indicating the insufficient representativeness of input data. They should be transparently reasoned and take the form of a supervisory decision. This preserves the integrity of the algorithmic system, while ensuring that the system remains flexible and adaptable to real circumstances and changes in risks. All this enhances the synergy between mathematical rigour and regulatory responsibility, which is the basis for sound and sustainable risk management in the context of preventing money laundering and terrorist financing.

8 Conclusion

Introducing a systematically defined and mathematically-based methodology for inherent risk assessment is a step forward in AML/CFT supervision. Using the method of minimising within-cluster variance to classify quantitative indicators allows for a high degree of objectivity, transparency and reproducibility of results, which is in line with contemporary regulatory trends. Application of this method allows for a fair and numerical risk categorisation, eliminating subjective oscillations that are frequent in manual scoring. At the same time, the model leaves room for necessary supervisory interventions, ensuring a balance between algorithmic precision and regulatory flexibility. Such an approach allows for the integration of empirical data and expert judgement, thus ensuring a holistic and adaptable risk assessment. The method of minimising within-cluster variance is one of many possible risk-valuation techniques, and its mathematical consistency, together with the possibility of automation in



programming languages such as Python or advanced Excel models, makes it extremely suitable for regulatory practice. This methodology can contribute to enhancing trust in supervisory systems, increasing operational efficiency and better channelling of supervisory resources towards actual risk hotspots. Ultimately, a combination of mathematical modelling and regulatory expertise turns into a key tool in addressing the complex and evolving challenges of financial crime at the EU level and beyond.

References

- Hardin, J., Quesada, L., Ye, L. and Horton, Nicolas J. (2025): *The Exchangeability Assumption for Permutation Tests of Multiple Regression Models*, Implications for Statistics and Data Science Educators.
- Chakraborty, G. (2022): *Cluster-Based Risk Assessment Using WCSS Minimization*, Journal of Quantitative Finance and Economics.
- Jia, K., Zhao, X. and Zhang, L. (2013): *Assessing Money Laundering Risk of Financial Institutions with AHP: Supervisory Perspective*, Journal of Financial Risk Management, Vol. 2, No. 1, 29-31.
- Tomić, B. (2010): *Uvod u statistiku i ekonometriku*, Zagreb, Školska knjiga.
- FATF (2021): *Guidance on Risk-Based Supervision*. Paris, FATF/OECD. Available at: <https://www.fatf-gafi.org/content/dam/fatf-gafi/guidance/Guidance-Risk-Based-Supervision.pdf.coredownload.inline.pdf>.
- Basel Committee on Banking Supervision (2005): *Studies on the Validation of Internal Rating Systems*, BIS.
- Basel Institute on Governance (2024): *Basel AML Index – Methodology*.
- Power, M. (2009): *The Risk Management of Everything*, Journal of Risk Finance.
- European Central Bank (2018): *Guide to Internal Models*.
- European Banking Authority (2020): *Guidelines on PD estimation, LGD estimation and the treatment of defaulted exposures*.

PUBLISHER

Croatian National Bank
Trg hrvatskih velikana 3
10000 Zagreb
T. +385 1 4564 555
www.hnb.hr

EDITOR-IN-CHIEF

Ljubinko Jankov

EDITORIAL BOARD

Alan Bobetko
Ana Vargek Stilinović
Boris Čota
Davor Galinec
Davor Kunovac
Evan Kraft
Ivo Bakota
Jurica Zrnc
Maja Bukovšak
Vedran Šošić

EXECUTIVE EDITOR

Katja Gattin Turkalj

TECHNICAL EDITOR

Nevena Jerak Muravec

The views expressed in this paper are not necessarily the views of the Croatian National Bank.

Those using data from this publication are requested to cite the source.

ISSN 1334-014X (online)

**Application of Optimised Discretisation Based on the Within-cluster Sum of Squares to
Assess the Inherent Risk of Money Laundering and Terrorist Financing**

ISSN 1334-014X (online)

