

ENTROPY TECHNIQUES FOR ROBUST MANAGEMENT DECISION MAKING IN HIGH-DIMENSIONAL DATA

Jan Kalina^{a,b}

^a*The Czech Academy of Sciences, Institute of Computer Science,
Pod Vodárenskou věží 2, 182 00 Prague 8, Czech Republic*

^b*Charles University, Faculty of Mathematics and Physics, Sokolovská 83,
186 75 Prague 8, Czech Republic*

(Received 15 January 2024; accepted 27 August 2024)

Abstract

Entropy, a key measure of chaos or diversity, has recently found intriguing applications in the realm of management science. Traditional entropy-based approaches for data analysis, however, prove inadequate when dealing with high-dimensional datasets. In this paper, a novel uncertainty coefficient based on entropy is proposed for categorical data, together with a pattern discovery method suitable for management applications. Furthermore, we present a robust fractal-inspired technique for estimating covariance matrices in multivariate data. The efficacy of this method is thoroughly examined using three real datasets with economic relevance. The results demonstrate the superior performance of our approach, even in scenarios involving a limited number of variables. This suggests that managerial decision-making processes should reflect the inherent fractal structure present in the given multivariate data. The work emphasizes the importance of considering fractal characteristics in managerial decision-making, thereby advancing the applicability and effectiveness of entropy-based methods in management science.

Keywords: information theory, high-dimensional data, uncertainty, robustness, management science

1. INTRODUCTION

Chaos theory serves as a paradigmatic framework elucidating the inherent chaos, randomness, uncertainty, and unpredictability of events in our

surroundings (Gleick, 2008). The mathematical underpinning of complex nonlinear dynamic systems reveals that specific chaotic systems manifest intricate fractal patterns characterized by their self-similarity and inherent self-development,

* Corresponding author: kalina@cs.cas.cz

exhibiting self-organization in various structures (Akhmet et al., 2020). The explication of complex and unpredictable systems falls within the purview of cybernetics, a discipline focused on elucidating feedback and control mechanisms.

Fractal patterns are not confined solely to the natural realm but can also be discerned in the societal framework, marked by an uneven and multifaceted evolution. The contemporary society is characterized by heightened instability, lack of control, unpredictability, and a complexity surpassing that of previous eras (Darian-Smith, 2022). The society's susceptibility to volatility, amplified by rapid technological advancements and further accentuated in the wake of the COVID-19 pandemic, is termed as its fractalization (Kalina, 2022).

Economic models face challenges in capturing the unpredictable cycles and chaotic trajectories of economic development, which exhibit fractal patterns with consequential implications for managerial forecasts (Slanina, 2013). Fractal-inspired methodologies have proven promising in financial management decision-making (Mosteanu, 2019) and portfolio management (Tilfani et al., 2020). However, their application in non-financial management tasks remains relatively underexplored. The incorporation of fractal perspectives holds potential for a more comprehensive understanding of managerial challenges beyond the financial sphere.

Entropy as a fundamental mathematical metric is widely acknowledged for its role in gauging the degree of chaos, diversity, and disorder within a specific system. While commonly employed in the scientific depiction of natural phenomena, it also finds application, albeit more metaphorically, in

characterizing societal dynamics. Although the precise evaluation of entropy for an entire society proves challenging, it is evident that a state of nature or a stable environment would exhibit a low entropy level. Conversely, the advent of globalization has been accompanied by a discernible rise in entropy. In tandem with the escalating global entropy, the complexity inherent in managerial decision-making has experienced a marked upsurge (Kahneman, 2011). The heightened levels of uncertainty often liken decision-making challenges to games of chance. Despite this, the strategic utilization of data remains a valuable asset in the marketplace and the ability to harness data effectively becomes increasingly crucial in navigating the intricacies of contemporary management (Himeur et al., 2023).

The complexity of available data, particularly in management contexts, is on the rise, owing to its heightened nonlinearity and large numbers of variables. In many instances, the data itself can be regarded as possessing a fractal structure characterized by a high fractal dimensionality (Jiřina & Jiřina, 2015). In such a case, statistical and machine learning methodologies should be deployed to transform raw data into valuable and interpretable knowledge, thereby reducing entropy by discerning meaningful distinctions between signal and noise (Reddy et al., 2020).

Addressing the challenges posed by high-dimensional data, dimensionality reduction techniques come into play. These methods work towards compressing information, eliminating redundant variables, and identifying only the most salient and predictable patterns (Chhikara et al., 2022). Traditional data analysis methods often prove inadequate for high-dimensional data, where a substantial number of variables are

at play. Consequently, there exists a gap in data analysis methodologies inspired by fractals and entropy, which could find application in multivariate data analysis, the construction of complex predictive models, optimization tasks, and multi-criterial decision-making processes, among others.

Section 2 provides an overview of entropy as a fundamental measure, emphasizing its recent applications within the domain of management science. In Section 3, a regularized version of the uncertainty coefficient, which is an entropy-based measure for categorical data, is proposed. The application of this coefficient is explored in Section 4, where it is posited as a valuable tool for pattern discovery, particularly in scenarios characterized by a multitude of categorical variables. In Section 5, the focus shifts to the presentation of a robust fractal-inspired methodology devised for estimating the covariance matrix of multivariate data. This novel method is subjected to scrutiny through its application to three multivariate datasets, demonstrating its efficacy and applicability. The results reveal the method not to be confined to high-dimensional datasets. In Section 6, conclusions are drawn. This structured progression serves to systematically unfold the contributions and applications of entropy, uncertainty coefficients, and fractal-inspired methodologies within the realm of management science.

2. ENTROPY IN MANAGEMENT SCIENCE

The (Shannon) entropy is defined for a discrete random variable X as

$$H(X) = -\sum p(x) \log p(x), \quad (1)$$

where the summation considers all the possible outcomes of X and $p(x)$ is the probability of an individual outcome x . Entropy and related characteristics, which are commonly inspired by thermodynamics, have been thoroughly studied in the field of information theory (Delgado-Bonal & Marshak, 2019). Evaluating entropy in management applications is often inaccessible or can be evaluated in a very specific situation (Czyż & Hauke, 2015). Much more useful tasks include exploiting entropy for evaluating related measures of uncertainty, such as the uncertainty coefficient studied later in Section 3. Also the Kullback-Leibler divergence (often denoted as relative entropy), which was used e.g. in the energy storage management application of Le et al. (2021), is derived from the concept of Shannon entropy. In this section, we recall and systematize recent management applications of entropy and search for interesting connections.

Maximizing entropy is natural in financial risk management and plays a key role in modern approaches to portfolio optimization. For example, Li & Zhang (2021) considered a maximization of a combination of variance and entropy in portfolio optimization. In a non-financial context, entropy maximization was used in the context of data fusion in Zamani et al. (2023) to integrate the results of multiple estimation methods in water quality management. The method is able to incorporate prior information by a Bayesian argument. Maximizing entropy was also used in Xiao et al. (2022) with a reinforcement learning aimed at finding an energy management strategy for an electric vehicle. There, energy efficiency was penalized by entropy in the process of modeling the interaction between the vehicle and the environment.

Entropy emerges as a widely recognized tool in the domain of Multi-Criteria Decision Making (MCDM) within management applications. A prevalent approach in MCDM simplifies the complex decision-making task by transforming it into a single-criterion decision-making problem. This is achieved through the creation of a weighted combination of all criteria, where the weight assigned to each criterion corresponds to its entropy, as evaluated across the training data. Such straightforward strategy turned out to be successful e.g. in Vujčić et al. (2017), Krstić & Fedajev (2020), and Fedajev et al. (2021). In all these papers, the method is denoted as the “entropy method”, which is unspecific and possibly misleading; we therefore suggest a renaming to “entropy-weighted single criterion decision making” (EW-SCDM). The principle of the method resembles constructing weights based on variability of available continuous variables; such approach is known as the inverse-variance weighting (Lee et al., 2016).

Punetha & Jain (2023) constructed a decision system for recommending the most suitable mobile phone for a given user; their analysis revealed the EW-SCDM method within an MCDM task to be outperformed by COPRAS, which is a more complex decision model abbreviating the complex proportional assessment of alternatives. Other complex decision models, for which a systematic comparison with the entropy method seems still missing, include TOPSIS, VIKOR, or ELECTRE (Pamučar et al., 2017).

The application of the max-entropy principle provides an alternative avenue for estimating unknown parameters, particularly in situations where direct evaluation of entropy may pose challenges. This approach involves considering the least favorable case, ensuring robustness in scenarios where

precise entropy assessment is impractical.

Illustrating this principle in the context of decision-making, Che et al. (2022) employed entropy maximization within the DEMATEL (Decision-Making Trial and Evaluation Laboratory) framework—a versatile method in multi-criterial decision-making that identifies complex causal relationships. In their study, the authors utilized the max-entropy principle to determine the initial influence matrix, thereby establishing the degree of direct influence between various factors (Šmidovnik & Grošelj, 2023). The max-entropy principle is analogous to considering least favorable distributions in statistics, where the situation that is the most difficult or least likely within a broad class of situations is considered as the most interesting one (Güney et al., 2021).

3. REGULARIZED UNCERTAINTY COEFFICIENT FOR CONTINGENCY TABLES

Managerial decision making is often shaped by uncertainties, and effective managers strive to anticipate and navigate these uncertainties (Love et al., 2022). In certain situations, it proves beneficial to make deliberate attempts to quantify uncertainty, even if only within specific, defined contexts (Lin et al., 2021). The uncertainty coefficient U derived from Shannon entropy represents a characteristic of uncertainty, which may be reliable only for data with a small number of variables. Our idea is to use regularization, which represents a common tool for extending data analysis tools to high-dimensional data (Kalina, 2024), and to propose a regularized version of U in this section. The novel version suitable (not only) for high-

dimensional data may play the role of an association measure between two categorical variables. Its possible application within pattern discovery will be suggested in Section 4.

Table 1. Contingency table $2 \times K$.

	$X=1$	$X=2$	...	$X=K$	Σ
$Y=1$	n_{11}	n_{12}	...	n_{1K}	$n_{1\cdot}$
$Y=0$	n_{21}	n_{22}	...	n_{2K}	$n_{2\cdot}$
Σ	$n_{\cdot 1}$	$n_{\cdot 2}$	...	$n_{\cdot K}$	n

Let us consider a $2 \times K$ contingency table shown as Table 1. We assume each column of the table to be generated by a binomial distribution. We understand X to be a categorical variable with K categories and Y to be a response variable and the unknown probability of $Y=1$ (success) in each column is denoted as π_k for $j=1, \dots, K$. The maximum likelihood estimator (MLE) of π_k is simply obtained as $\hat{\pi}_k = n_{1k}/n_{\cdot k}$. In this notation, we can say that the data are observed in two different groups (supervised situation) corresponding to the two levels of the response. We propose now to consider regularized estimators of π_k denoted as π_k^* in the form

$$\pi_k^* = (1 - \lambda) \frac{n_{1k}}{n_{\cdot k}} + \lambda \frac{n_{1\cdot}}{n}, \quad k = 1, \dots, K, \quad (2)$$

for a given $\lambda \in [0, 1]$. This estimator for the probability of success in the j -th group is based on π_k shrunk to the overall MLE across categories. The estimator (2) for the k -th category borrows information from all the remaining categories as it is common for regularization for continuous data, e.g. for covariance matrix estimates for high-dimensional data (Sadik et al., 2023).

The regularization in (2) represents a

penalization of $\hat{\pi}_k$, which is meaningful when the columns of the table correspond e.g. to different strata especially for a large K . In such situations, regularized approaches may be preferable due to the bias-variance-tradeoff, i.e. individual estimation tools may much improve the variability at the cost of increasing their bias. It may not be appropriate in some other designs, e.g. if there is a natural ordering of the columns of the contingency tables. The estimate (2) is a compromise between $n_{1k}/n_{\cdot k}$ and $n_{1\cdot}/n$; $\lambda=0$ clearly corresponds to the plain MLE.

The biased estimates (2) may be plugged into standard formulas for various association measures for Table 1. This allows e.g. to formulate a regularized version of Pearson's χ^2 statistic as well as of the likelihood ratio statistic G^2 , which both are well known statistics of the test of independence for contingency tables. These test statistics (without evaluating their p-values) may also be exploited as the basis for obtaining regularized distance measures for categorical data.

The uncertainty coefficient U , which is also known as Theil's index U , represents a common measure of accuracy of classification methods for categorical data; see p. 57 of Agresti (2002). It represents an (asymmetric) association measure between Y and X . The population version is formally defined as

$$U(Y|X) = \frac{H(Y)H(Y|X)}{H(Y)} = \frac{H(X) + H(Y) - H(X, Y)}{H(Y)} = \frac{I(X, Y)}{H(Y)}, \quad (3)$$

where $H(X)$ denotes the entropy of the categorical variable X , $H(X, Y)$ denotes the joint entropy of variables X and Y ; $H(Y|X)$

denotes the conditional entropy, and $I(X, Y)$ denotes the joint information of X and Y . The equivalences in (3) use the fact that $H(Y|X) = H(X, Y) - H(X)$. The coefficient (3) may be perceived as a conditional entropy coefficient of the response. It evaluates the proportional reduction in entropy, i.e. the contribution of the sum $H(X) + H(Y)$ compared to $H(X, Y)$ relatively to $H(X)$, in other words a conditional contribution of X to the knowledge of Y . It holds that $U(Y|X) \in [0, 1]$ and particularly $U(Y|X) = 0$ in case of independence of X and Y , i.e. in case that X does not carry any information about Y .

Let us now propose a novel regularized version of the uncertainty coefficient, which is obtained by regularizing an empirical version of (3). Let us define the regularized version as

$$U^*(Y|X) = \frac{-\sum_{i=1}^2 \frac{n_{i\cdot}}{n} - \sum_{k=1}^K \pi_k^* \log \pi_k^* + \sum_{i=1}^2 \sum_{k=1}^K \frac{n_{ik}}{n} \log \frac{n_{ik}}{n}}{-\sum_{k=1}^K \pi_k^* \log \pi_k^*} \quad (4)$$

with $\pi_1^*, \dots, \pi_K^*$ defined in (2); the definition is meaningful for data with $n_{1\cdot} > 0$ and $n_{2\cdot} > 0$. Here, $\log$ denotes the natural logarithm. Clearly, $U^*(Y|X) \in [0, 1]$.

The coefficient $U^*(Y|X)$ depends on λ through (4). It is even possible to obtain an explicit expression for the optimal value of λ as an adaption of the general result of Ledoit & Wolf (2022). The optimal value, denoted here as $\lambda^\dagger$, is derived as the asymptotically optimal value minimizing the mean square error among all values of (2) over all $\lambda \in [0, 1]$ for $n \rightarrow \infty$. This asymptotically optimal value of λ has the form

$$\lambda^\dagger = \max \left\{ 0, \frac{1 - \sum_{k=1}^K \left(\frac{n_{1k}}{n_{\cdot k}} \right)^2}{(n-1) \sum_{k=1}^K \left(\frac{n_{1\cdot}}{n} - \frac{n_{1k}}{n_{\cdot k}} \right)^2} \right\}. \quad (5)$$

If the probabilities $\pi_1, \dots, \pi_K$ are close to homogeneous across all K groups, the value of $\lambda^\dagger$ is likely to be close to 0, which corresponds to negligible (or no) effect of regularization. This corresponds to intuition, because the regularization is desirable especially in a very heterogeneous situation with very diverse results (i.e. due to the presence of rare events) across individual columns of Table 1. The explicit expression for $\lambda^\dagger$ may simplify various complications in practical tasks; one such method for a pattern discovery method for categorical data is suggested in Section 4.

4. PATTERN DISCOVERY FOR CATEGORICAL DATA

Pattern discovery as the identification of patterns, characterized as the exploration for key distinctions between two groups of data samples, constitutes a crucial tool of exploratory data analysis (EDA) (Subramanian et al., 2020). Pattern discovery for continuous dynamic data was used e.g. in portfolio optimization in Martins & Neves (2020). In this section, a possible approach for pattern discovery for categorical data is suggested, exploiting the regularized uncertainty coefficient (4) in the task to find variables contributing the most to the discrimination between two given groups. The method, which is intended as a dimensionality reduction method (not only) for a large number of categorical variables, will be described on the example of a three-dimensional contingency table of size $2 \times 2 \times K$, while it is designed to work also for categorical data with a much larger number of variables.

Table 2 is assumed to follow a multinomial model, where n is the total

Table 2. Contingency table $2 \times 2 \times K$

	$X = 1$		$X = 2$			$X = K$	
	$Z = 1$	$Z = 0$	$Z = 1$	$Z = 0$	...	$Z = 1$	$Z = 0$
$Y = 1$	n_{111}	n_{121}	n_{112}	n_{122}	...	n_{11K}	n_{12K}
$Y = 0$	n_{211}	n_{221}	n_{212}	n_{222}	...	n_{21K}	n_{22K}

number of counts. We consider the binomial model in each of the total number of $2K$ columns of Table 2. From the point of view of a classification task, we understand Y to represent a response variable. We assume K to be large, while we are interested in classification to two groups based on two levels of the variable X . The aim is to find relevant patterns, i.e. to detect the variables that are strong predictors for the response. Various approaches to pattern discovery have been available for the analysis of categorical data; commonly, they exploit the Pearson's χ^2 test statistic based on adjusted residuals (Agresti, 2002). Our novel approach allows to analyze complicated quantitative and qualitative associations among categorical random variables and to discover patterns in a dataset and thus to generate interesting hypotheses from data. The subsequent analysis may be based only on the detected variables, i.e. the procedure may be exploited for dimensionality reduction purposes.

The idea is to consider all possible marginal tables obtained from Table 2 by ignoring the effect of some of the categorical variables. The regularized uncertainty coefficient will be evaluated for all such tables. In this way, such approach reveals simple associations among variables, i.e. low-order patterns (patterns with a small dimension). The regularized uncertainty coefficient may be computed for various combinations of the available variables. All such combinations should clearly retain Y in

order to evaluate the contribution of other variables to its knowledge. The coefficient (4) together with the optimal value of the shrinkage intensity $\lambda^{\dagger}$ (5) quantifies in each situation the conditional contribution of the considered variables to the knowledge of Y as explained in Section 3.

Table 3. An auxiliary contingency table considered within the procedure of Section 4

	$X = 1$	$X = 2$	...	$X = K$
$Y = 1$	$n_{1 \cdot 1}$	$n_{1 \cdot 2}$	...	$n_{1 \cdot K}$
$Y = 0$	$n_{2 \cdot 1}$	$n_{2 \cdot 2}$	...	$n_{2 \cdot K}$

To illustrate low-order patterns obtained from Table 2, let us have a look on the particular table shown as Table 3, which is one of the $2 \times K$ tables considered within the procedure and evaluated by methods suitable for (2). In this way, various tables will be analyzed comparing different patterns in the data. Only the patterns are considered to be detected and recommended for a subsequent analysis that lead to large values of the coefficient $U^*(Y|X)$. Two possible approaches could be used: (i) selecting a fixed number (for example 10) of the most relevant patterns, or (ii) selecting the patterns that lead to $U^*(Y|X)$ exceeding a given threshold. The property $U^*(Y|X) \in [0, 1]$ suggests to choose a value closer to 1 such as 0.7 or 0.8, although there has been no agreement about a single value of the uncertainty coefficient, which would be considered acceptably or significantly high (Vourdas, 2020).

On the whole, the pattern discovery of this section may be performed for categorical data with a much higher number of variables than in Table 2, still perceiving the data as observed in two different groups. The method is designed tailor-made to have a strong ability to detect the patterns that contribute the most to the difference between the two groups.

Keziou & Regnault (2017) investigated estimates of mutual information based on entropy-based measures such as f-divergences in a semiparametric context; the estimates were to derive independence tests with optimality properties. The semiparametric ideas were extended by Broniatowski (2021) to derive minimum divergence estimators as a broad class of methods, which generalize e.g. the Pearson's χ^2 statistic. The approach has robustness properties and could be also exploited within the pattern discovery of this section.

5. ROBUST COVARIANCE MATRIX WITH RECIPROCAL WEIGHTS FOR HIGH-DIMENSIONAL DATA

Various statistical tasks important for managers require to estimate the covariance matrix of continuous multivariate data. To give some examples, Li et al. (2023) exploited covariance matrices within clustering applied to developing an effective energy management strategy. Boas et al. (2023) used a covariance matrix adaptation evolution strategy (CMA-ES) for the quality management of human-robot collaboration. Wielicka-Gańczarczyk & Jonek-Kowalska (2023) estimated the covariance matrix as a variability characteristic within a model for smart city management. Covariance matrices were used within a classification task solved

in a credit risk management context in Kalina (2022), who did not however consider fractal-inspired reciprocal weights. Here, a robust estimator of the covariance matrix, which is tailor-made for high-dimensional data, is proposed exploiting reciprocal weights assigned to individual observations.

The fractal-inspired weights in the form of a reciprocal weight function were used in a classification context by Jiřina & Jiřina (2015), who associated reciprocal weights with fractals through the Zips's law; this assumption (rather than law) states that ranks of distances for pairs of observations are very often distributed according to the Zipf's distributions for real data.

The MRWCD estimator was proposed very recently in Kalina & Tichavský (2022) as a highly robust tool for multivariate contaminated data. It abbreviates the minimum regularized weighted covariance determinant and is based on assigning a weight to every observation. As a novelty, we now propose to use the MRWCD estimator with the reciprocal weights using the weight function

$$\psi(t) = \frac{1}{t}, \quad t \in (0,1). \quad (6)$$

Our numerical experiments aim at comparing robust multivariate estimators over 3 real datasets all coming from publicly available repositories. Particularly, they reveal the performance of an implicitly weighted MRWCD estimator of Σ with fractal-inspired weights. The datasets Aircraft and Delivery come from Rousseeuw & Leroy (1987); the Housing dataset comes from the UCI public repository. These three datasets, which are known as benchmarking data for robust estimation, contain no missing values. The housing dataset is the

largest and has acquired attention in the field of urban planning (Chen et al., 2023); the other two smaller datasets are used for comparisons.

To estimate the mean and covariance matrix of the data, we use the following estimators: MLE (maximum likelihood, i.e. mean and empirical covariance matrix), MCD (Rousseeuw & Leroy, 1987) with the trimming constant $3/4$, MM-estimators (Tatsuoka & Tyler, 2000), and MRWCD with reciprocal weights. For each of the estimators, we evaluate a summary measure of variability denoted here as $\mathcal{M}$, which sums the variance estimates across all (say p) individual variables. For a given estimate $\hat{\Sigma}$ of the true covariance matrix, the measure $\mathcal{M}$ is defined by

$$\mathcal{M} = \sum_{j=1}^p \widehat{var} \hat{\Sigma}_{jj}, \quad (7)$$

where $\widehat{var}$ is the nonparametric bootstrap estimate. In other words, $\mathcal{M}$ adds up estimates of diagonal elements of the covariance matrix of the variances of the given data. It is desirable to use estimators that achieve a small value of $\mathcal{M}$ across various realistic datasets. Principles of nonparametric bootstrap have been many times justified for the context of covariance matrix estimation; a recent application was

presented e.g. in Al-Hadeethi et al. (2022).

The results are presented in Table 4, which also gives the number of observations n and number of variables p of each dataset. Because we can expect the results to depend not only on outliers but also on the correlation structure of the data, the tables include $\det(S)$ as the determinant of the empirical covariance matrix S and κ as the condition number of S . MLE, which is the only non-robust method here, performs quite well for the Aircraft and Housing datasets, which do not contain severe outliers. The Delivery dataset is the contrary and only MM and MRWCD perform well there.

MCD and MM are robust but not regularized, i.e. unsuitable if the covariance matrix is ill-conditioned. MRWCD, which is robust and regularized, is suitable also for ill-conditioned data and outperforms the remaining estimators for all the selected datasets. MM-estimators represent the golden standard and stays somewhat behind MRWCD. This is true for datasets, which have a high κ and thus problems with numerical stability of the covariance matrix, as well as for those with a reasonably low κ . MCD is the worst in all the three datasets, which confirms the findings of Kalina & Tichavský (2022).

Table 4. Analysis of the three datasets of Section 5. The value of Mis reported for 4 different estimators

		Dataset		
		Aircraft	Delivery	Housing
n		23	25	506
p		4	2	11
$\det(S)$		0.06	0.32	0.0005
κ		40.4	11.2	21.3
Value of $\mathcal{M}$	MLE	0.67	0.35	0.06
	MCD	7.79	0.41	0.17
	MM	1.96	0.11	0.14
	MRWCD	0.60	0.10	0.03

6. CONCLUSIONS

Despite the growing significance of disorder in the society, the indispensability of data-driven methods in management persists. Still, approaches influenced by fractals prove to be more pertinent than conventional tools. Consequently, managerial decision making grounded in available data should mirror the fractal structure inherent in multivariate data. This necessitates managers to grasp fractals and employ more sophisticated methods compared to simplistic approaches like e.g. the entropy method, referred to here as EW-SCDM, within MCDM tasks.

The experiments presented in Section 5 demonstrate the advantageous performance of the novel version of the MRWCD estimator, even for data with a small number of variables. The inadequacy of standard covariance matrix estimators becomes evident in high-dimensional datasets. It is essential to note that the analysis is based on only three selected datasets. Future research should prioritize exploring alternative approaches tailored for high-dimensional data, be it categorical or continuous, with practical applications in management.

While this paper does not delve into the realm of fractal management within an organization possessing a fractal structure, the intricate relationships between the two warrant a dedicated research paper (Rezazadeh et al., 2023). It is important to acknowledge that entropy within an organization can be advantageous, yet in certain contexts, proactive measures may be more effective in mitigating its effects.

Acknowledgment

The work was supported by the project 24-11146S of the Czech Science Foundation.

References

- Agresti, A. (2002). *Categorical data analysis*. 2nd edn. Hoboken, NJ: Wiley.
- Akhmet, M., Fen, M.O., & Alejaily, E.M. (2020). *Dynamics with chaos and fractals*. Cham, Switzerland: Springer.
- Al-Hadeethi, H., Abdulla, S., Diykh, M., Deo, R.C., & Green, J.H. (2022) An eigenvalue-based covariance matrix bootstrap model integrated with support vector machines for multichannel EEG signals analysis, *Front. Neuroinform.* 15, 808339.
- Boas, A.V., André, J., Cerqueira, S.M., & Santos, C.P. (2023). A DMPs-based approach for human-robot collaboration task quality management. *IEEE International Conference on Autonomous Robust Systems and Competitions ICARSC 2023*, 226–231.
- Broniatowski, M. (2021). Minimum divergence estimators, maximum likelihood and the generalized bootstrap. *Entropy*, 23, 185.
- Che, Y., Deng, Y., & Yuan, Y.H. (2022). Maximum-entropy-based decision-making trial and evaluation laboratory and its application in emergency management. *Journal of Organizational and End User Computing*, 34, 1–16.
- Chen, Y., Jiao, J., & Farahi, A. (2023). Disparities in affecting factors of housing price: A machine learning approach to the effects of housing status, public transit, and density factors on single-family housing

ЕНТРОПИЈСКЕ ТЕХНИКЕ ЗА ПОУЗДАНО ДОНОШЕЊЕ МЕНАѢЕРСКИХ ОДЛУКА У УСЛОВИМА ВИШЕДИМЕНЗИОНАЛНИХ ПОДАТАКА

Jan Kalina

Извод

Ентропија, као кључна мера хаоса или разноликости, последњих година налази све шире примене у науци о менаѢменту. Ипак, традиционални приступи засновани на ентропији показују ограничену ефикасност када је реч о анализи вишедимензионалних скупова података. У овом раду се предлаже нови коефицијент неизвесности, заснован на ентропији, који је прилагођен категоријским подацима, као и метода за откривање образаца погодна за примену у менаѢерским ситуацијама. Поред тога, представљена је поуздана техника инспирисана фракталима за процену коваријантних матрица у мултиваријатним подацима. Ефикасност ове методе детаљно је анализирана кроз три скупа података са економском релевантношћу. Резултати потврђују супериорне перформансе предложеног приступа чак и у сценаријима са ограниченим бројем променљивих. Ова истраживања указују на потребу да се у процесима доношења менаѢерских одлука узму у обзир урођене фракталне структуре присутне у вишедимензионалним подацима. Рад наглашава значај разматрања фракталних карактеристика у менаѢерским одлукама, чиме се унапређује применљивост и ефикасност ентропијских метода у науци о менаѢменту.

Кључне речи: Теорија информација, вишедимензионални подаци, неизвесност, робусност, наука о менаѢменту.

price. Cities, 140, 104432.

Chhikara, P. Jain, N., Tekchandani, R., & Kumar, N. (2022). Data dimensionality reduction techniques for Industry 4.0: Research results, challenges, and future research directions. *Software Practice and Experience*, 52, 658–688.

Czyż, T. & Hauke, J. (2015). Entropy in regional analysis. *Quaestiones Geographicae*, 34, 69–78.

Darian-Smith, E. (2022). *Global burning: Rising antidemocracy and the climate crisis*. Stanford, CA: Stanford University Press.

Delgado-Bonal, A. & Marshak, A. (2019). Approximate entropy and sample entropy: A comprehensive tutorial. *Entropy*, 21, 541.

Fedajev, A., Radulescu, M., Babucea, A.G., Mihajlovic, V. (2021). Real convergence in EU: Is there a difference between the effects of the pandemic and the global economic crisis? *Politická ekonomie*, 69, 571–594.

Gleick, J. (2008). *Chaos: Making a new science*. NY: Viking Press.

Güney, Y., Tuac, Y., Özdemir, S., & Arslan, O. (2021). Robust estimation and variable selection in heteroscedastic regression model using least favorable distribution. *Computational Statistics*, 36, 805–827.

Himeur, Y., Elnour, M., Fadli, F., Meskin, N., Petri, I., et al. (2023). AI-big data analytics for building automation and

- management systems: A survey, actual challenges and future perspectives. *Artificial Intelligence Review*, 56, 4929–5021.
- Jiřina, M. & Jiřina, M. (2015). Classification using the Zipfian kernel. *Journal of Classification*, 32, 305–326.
- Kahneman, D. (2011). *Thinking, fast and slow*. NY: Farrar, Straus and Giroux.
- Kalina, J. (2024). Regularized least weighted squares estimator in linear regression. *Communications in Statistics–Simulation and Computation*. Accepted.
- Kalina, J. (2022). Decision making reflecting the fractalization of the society. *Serbian Journal of Management*, 17, 207–218.
- Kalina, J. & Tichavský, J. (2022). The minimum weighted covariance determinant estimator for high-dimensional data. *Advances in Data Analysis and Classification*, 16, 977–999.
- Keziou, A. & Regnault, P. (2017). Semiparametric estimation of mutual information and related criteria: Optimal test of independence. *IEEE Transactions on Information Theory*, 63, 57–71.
- Krstić, S & Fedajev, A. (2020). The role and importance of large companies in the economy of the Republic of Serbia. *Serbian Journal of Management*, 15, 335–352.
- Le, J., Liao, X., Zhang, L., & Mao, T. (2021). Distributionally robust chance constrained planning model for energy storage plants based on Kullback-Leibler divergence. *Energy Reports*, 7, 5203–5213.
- Ledoit, O. & Wolf, M. (2022). The power of (non-)linear shrinking: A review and guide to covariance matrix estimation. *Journal of Financial Econometrics*, 20, 187–218.
- Lee, C.H., Cook, S., Lee, J.S., & Han, B. (2016). Comparison of two meta-analysis methods: Inverse-variance-weighted average and weighted sum of Z-scores. *Genomics & Informatics*, 14, 173–180.
- Li, B. & Zhang, R. (2021). A new mean-variance-entropy model for uncertain portfolio optimization with liquidity and diversification. *Chaos, Solitons, & Fractals*, 146, 110842.
- Li, M., Wang, Y., Yu, P., Sun, Z., & Chen, Z. (2023). Online adaptive energy management strategy for fuel cell hybrid vehicles based on improved cluster and regression learner. *Energy Conversion and Management*, 292, 117388.
- Lin, L., Bao, H., & Dinh, N. (2021). Uncertainty quantification and software risk analysis for digital twins in the nearly autonomous management and control systems: A review. *Annals of Nuclear Energy*, 160, 108362.
- Love, P.E.D., Ika, L.A., & Pinto, J.K. (2022). Homo heuristicus: From risk management to managing uncertainty in large-scale infrastructure projects. *IEEE Transactions on Engineering Management*, 71, 1940–1949.
- Martins, T.M. & Neves, R.F. (2020). Applying genetic algorithms with speciation for optimization of grid template detection in financial markets. *Expert Systems with Applications*, 147, 113191.
- Mosteanu, N.R. (2019). Intelligent tool to prevent economic crisis–fractals. A possible solution to assess the management of financial risk. *Quality–Access to Success*, 20, 13–17.
- Pamučar, D.S., Božanić, D., & Ranđelović, A. (2017). Multi-criteria decision making: An example of sensitivity analysis. *Serbian Journal of Management*, 12, 1–27.
- Punetha, N. & Jain, G. (2023). Integrated Shannon entropy and COPRAS optimal

- model-based recommendation network. *Evolutionary Intelligence*, 17 (1), 1-13
- Reddy, G.T., Reddy, M.P.K., Lakshmana, K., Kaluri, R., Rajput, D.S., et al. (2020). Analysis of dimensionality reduction techniques on Big Data. *IEEE Access*, 8, 54776–54788.
- Rezazadeh, F., Rezazadeh, S., & Rezazadeh, M. (2023). Fractal organizations and employee-organization relationship dynamics. In Faghih, N. (ed.), *Time and Fractals. Contributions to Management Science*. Cham, Switzerland: Springer.
- Rousseeuw, P.J. & Leroy, A.M. (1987). *Robust regression and outlier detection*. NY: Wiley.
- Sadik, S., Et-tolba, M., & Nsiri, B. (2023). A novel regularization-based optimization approach to sparse mean-reverting portfolios selection. *Optimization and Engineering*, 24, 2549-2577.
- Slanina, F. (2013). *Essentials of econophysics modelling*. Oxford, England: Oxford University Press.
- Subramanian, I., Verma, S., Kumar, S., Jere, A., Anamika, K. (2020). Multi-omics data integration, interpretation, and its application. *Bioinformatics and Biology Insights*, 2020, 14, 1177932219899051.
- Šmidovnik, T. & Grošelj, P. (2023). Solution for convergence problem in DEMATEL method: DEMATEL of finite sum of influences. *Symmetry*, 15 (7), 1357.
- Tatsuoka, K.S., & Tyler, D.E. (2000). On the uniqueness of S-functionals and M-functionals under nonelliptical distributions. *Annals of Statistics*, 28, 1219–1243.
- Tilfani, O., Ferreira, P., & El Boukfaoui, M.Y. (2020). Multiscale optimal portfolios using CAPM fractal regression: Estimation for emerging stock markets. *Post-Communist Economies*, 32, 77–112.
- Vourdas, A. (2020). Uncertainty relations in terms of the Gini index for finite quantum systems. *Europhysics Letters*, 130, 20003.
- Vujčić, M.D., Papić, M.Z., & Blajević, M.D. (2017). Comparative analysis of objective techniques for criteria weighing in two MCDM methods on example of an air conditioner selection. *Tehnika*, 67, 422–429.
- Wielicka-Gańczarczyk, K. & Jonek-Kowalska, I. (2023). Perceptions and attitudes toward risks of city administration employees in the context of smart city management. *Smart Cities*, 6, 1325–1344.
- Xiao, B., Yang, W., Wu, J., Walker, P.D., & Zhang, N. (2022). Energy management strategy via maximum entropy reinforcement learning for an extended range logistic vehicle. *Energy*, 253, 124105.
- Zamani, M.G., Nikoo, M.R., Niknazar, F., Al-Rawas, G., Al-Wardy, M., & Gandomi, A.H. (2023). A multi-model data fusion methodology for reservoir water quality based on machine learning algorithms and Bayesian maximum entropy. *Journal of Cleaner Production*, 416, 137885.