

國立臺灣大學電機資訊學院電機工程學系

碩士論文

Department of Electrical Engineering

College of Electrical Engineering and Computer Science

National Taiwan University

Master's Thesis



二維離散樣本熵：一種穩定且快速的紋理分析方法

Two-dimensional Dispersion Sample Entropy: A Robust
and Faster Method to Analyze Image Textures

陳英凡

Ying-Fan Chen

指導教授：顏嗣鈞 博士

Advisor: Hsu-Chun Yen, Ph.D.

中華民國 113 年 8 月

August, 2024

國立臺灣大學碩士學位論文
口試委員會審定書
MASTER'S THESIS ACCEPTANCE CERTIFICATE
NATIONAL TAIWAN UNIVERSITY

二維離散樣本熵: 一種穩定且快速的紋理分析方法
Two-dimensional Dispersion Sample Entropy: A Robust and Faster Method to
Analyze Image Textures

本論文係陳英凡 (R11921111) 在國立臺灣大學電機工程學系完成之碩士學位論文，於民國 113 年 7 月 23 日承下列考試委員審查通過及口試及格，特此證明。

The undersigned, appointed by the Department of Electrical Engineering on 23 July 2024 have examined a Master's thesis entitled above presented by Ying-Fan Chen (R11921111) candidate and hereby certify that it is worthy of acceptance.

口試委員 Oral examination committee:

顏尚穎

(指導教授 Advisor)

雷欽陞

郭斯齊

王柏堯

系主任 Director:

李建模



Acknowledgements

研究所的時光過得很快，但這兩年卻讓我成長了許多，不僅僅是在專業技能上的進步，也在心態上更為穩健。

在這兩年的時光裡，首先非常感謝顏嗣鈞老師的幫助與指導，不管是在論文還是課業上都讓我獲益良多。實驗室自由開明的風氣也是我非常喜歡的，也很感謝兩年來實驗室的同學和學長姐給予的建議與討論。

此外，也非常感謝在微軟實習遇到的同事和正職，在軟體開發，求職和人生規劃上都給予我很多實質的建議與幫助。

最後，感謝父母不管是對於我讀研究所的決定，還是選擇轉換專業領域的豪賭，給予的支持與鼓勵。只能說一路走來，雖然我一個人的力量渺小，但總是能收到許多人的幫助，對此我非常感激，也感謝和祝福每一個幫助過我的人都能遇到屬於自己的貴人。



摘要

信息理論已在各個領域廣泛應用，產生了用於分析時間序列數據的基於熵的度量。本文關注將這些度量擴展到更高維度，對於像圖像分析和分類這樣的任務至關重要。雖然二維樣本熵（SampEn_{2D}）在生物醫學圖像分析中顯示了潛力，但其計算效率低下且無法處理大尺寸圖像的限制限制了其實用性。相反，二維離散熵（DispEn_{2D}）提供了一種更快的替代方案，但在分類具有高不規則性的紋理方面存在困難。為了解決這些限制，本研究介紹了二維離散樣本熵（DispSampEn_{2D}），充分利用了二維樣本熵和二維離散熵的優勢，同時彌補了它們的缺陷。通過使用合成圖像和真實紋理數據集進行驗證，展示了二維離散樣本熵在分類任務中的有效性。此外，本研究強調了利用不同嵌入維度的熵向量的優勢，由於它們的互補特性，導致了改善的分類準確性。

關鍵字：資訊理論、紋理分析、樣本熵、離散熵、影像處理



Abstract

Information theory has found wide applications across diverse domains, giving rise to entropy-based metrics for analyzing time-series data. This thesis focuses on extending these metrics to higher dimensions, a necessity for tasks like image analysis and classification. While two-dimensional sample entropy (SampEn_{2D}) holds promise in biomedical image analysis, its computational inefficiency and inability to handle large-size images limit its utility. Conversely, two-dimensional dispersion entropy (DispEn_{2D}) offers a quicker alternative but struggles with classifying textures exhibiting high irregularity. To overcome these limitations, this study introduces two-dimensional dispersion sample entropy (DispSampEn_{2D}), leveraging the strengths of SampEn_{2D} and DispEn_{2D} while mitigating their weaknesses. Experiments using synthetic images and real-world texture datasets demonstrate the effectiveness of DispSampEn_{2D} in classification tasks. Furthermore, this research underscores the advantages of utilizing entropy vectors with diverse embedding dimensions, resulting in improved classification accuracy due to their comple-

mentary characteristics.

Keywords: Information theory, Texture analysis, Sample entropy, Dispersion entropy, Image processing

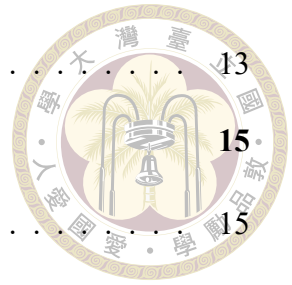




Contents

	Page
Acknowledgements	i
摘要	ii
Abstract	iii
Contents	v
List of Figures	vii
List of Tables	ix
Chapter 1 Introduction	1
Chapter 2 Related Work	3
2.1 Two-dimensional Dispersion Entropy	3
2.2 Two-dimensional Sample Entropy	5
Chapter 3	
Two Dimensional	
Dispersion Sample Entropy	8
3.1 Methodology	8
3.2 Suggested Value for Dispersion Classes	9
3.3 Linear Approach for Computation	10
Chapter 4 Datasets	12
4.1 Synthetic Datasets	12

4.2	Real-world Datasets	43
Chapter 5	Experiments and Results	15
5.1	Validation of DispSampEn_{2D}	15
5.2	Performance on Brodatz dataset	17
5.3	Performance on Kylberg dataset	19
5.4	Noise Resistance	21
5.5	Entropy Vector	25
5.6	Computation Time	28
Chapter 6	Conclusion	32
	References	34

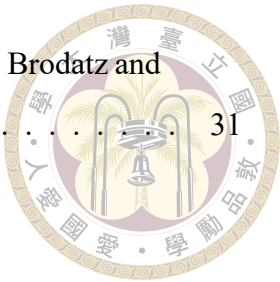




List of Figures

2.1	Example of first and last embedding vectors for SampEn _{2D}	6
4.1	Results of MIX _{2D} with different p	13
4.2	Samples from each group used in the normalized Brodatz dataset.	14
4.3	Samples from each group used in the Kylberg dataset.	14
4.4	MIX _{2D} applied on D5 with varying p	14
4.5	MIX _{2D} applied on blanket1 with varying p	14
5.1	Entropy of different metrics plotted against the value p of the MIX process.	16
5.2	Entropy for each group in the Brodatz dataset.	18
5.3	Entropy for each group in the Kylberg dataset.	20
5.4	Accuracy of classifying the Brodatz dataset with varying p for different metrics, with $m = 2$	21
5.5	Accuracy of classifying the Kylberg dataset with varying p for different metrics, with $m = 2$	22
5.6	DispEn _{2D} for each group from the Brodatz dataset.	23
5.7	DispEn _{2D} for each group from the Kylberg dataset.	24
5.8	Accuracy on real-world datasets, using entropy vectors with varying embedding dimensions as classifier input.	26
5.9	Confusion matrices for using different entropy vectors on the Brodatz dataset.	29
5.10	Confusion matrices for using different entropy vectors on the Kylberg dataset.	30

5.11 Accuracy per group for using different entropy vectors on the Brodatz and
Kylberg datasets. 31





List of Tables

5.1	Accuracy of classifying the Brodatz dataset for different metrics and embedding dimensions.	19
5.2	Accuracy of classifying the Kylberg dataset for different metrics and embedding dimensions.	21
5.3	Accuracy and respective undefined value ratios for SampEn _{2D} on the noisy Brodatz datasets.	25
5.4	Accuracy and respective undefined value ratios for SampEn _{2D} on the noisy Kylberg datasets.	25
5.5	Computation time required for calculating the Brodatz dataset for different metrics.	28
5.6	Computation time required for calculating the Kylberg dataset for different metrics.	28



Chapter 1 Introduction

Information theory has been applied to several fields, recently, the concept of entropy has been introduced to analyze the regularity and complexity on the time-series data. Multiple entropy metrics have been proposed, including distribution entropy (DistrEn_{1D}) [1], permutation entropy (PerEn_{1D}) [2] and dispersion entropy (DispEn_{1D}) [3] which are derived from Shannon entropy. Others such as approximate entropy (ApEn_{1D}) [4], sample entropy (SampEn_{1D}) [5, 6] and fuzzy entropy (FuzzEn_{1D}) [7] are based on the concept of the conditional entropy. In the past decade, these entropy metrics have been extended to higher dimensions, particularly for analyzing image textures and facilitating classification tasks in two-dimensional space. [8–17].

Two-dimensional sample entropy (SampEn_{2D}) [10, 11] has emerged as a potent tool for feature extraction from images, revealing the underlying regularity of patterns. It has been demonstrated that SampEn_{2D} is particularly useful in biomedical image analysis, serving as a texture descriptor to distinguish sural nerve images across different age groups of rats [10, 11]. However, the shortcomings of SampEn_{2D} are apparent, as the disadvantages inherent from SampEn_{1D} are magnified in its two-dimensional counterpart. Calculating SampEn_{2D} for large-size images is inefficient, primarily due to the $O(n^2)$ computational complexity of the naive sample entropy algorithm. Furthermore, larger embedding dimensions increase the likelihood of SampEn_{2D} yielding undefined values.

In comparison, two-dimensional dispersion entropy (DispEn_{2D}) [14] offers a faster and more stable alternative. By essentially computing a form of Shannon entropy, replacing probability with the occurrence frequency of patterns, DispEn_{2D} avoids the issue of undefined values. Nonetheless, DispEn_{2D} demonstrates limited efficacy in differentiating textures with higher irregularity, exhibiting poor noise resistance in classifying real-world texture datasets.

This thesis presents two main contributions. Firstly, a novel entropy metric termed two-dimensional dispersion sample entropy (DispSampEn_{2D}) is introduced. Derived from the concept of conditional entropy similar to SampEn_{2D} , DispSampEn_{2D} offers an implementation in linear time, making it considerably faster than SampEn_{2D} in computation. Additionally, DispSampEn_{2D} ensures the avoidance of undefined values under specific conditions and demonstrates improved robustness to noise compared to DispEn_{2D} and SampEn_{2D} . The validation of this new metric is conducted using synthesized images with mixed processes, and its performance is evaluated through classification tasks on two real-world texture datasets: the normalized Brodatz dataset and the Kylberg dataset. Notably, DispSampEn_{2D} can also be applied in one-dimensional space while preserving all its properties. Secondly, it is demonstrated that an entropy vector formed with different embedding dimensions yields enhanced accuracy in classification tasks due to the complementary nature of different embedding dimensions.



Chapter 2 Related Work

In this chapter, the fundamental algorithms of DispEn_{2D} and SampEn_{2D} are introduced. These metrics will be compared with our proposed DispSampEn_{2D} in Chapter 5. Furthermore, a comprehensive understanding of these two metrics can offer valuable insight into DispSampEn_{2D} .

2.1 Two-dimensional Dispersion Entropy

DispEn_{2D} extends DispEn_{1D} into two-dimensional space. Assuming an input image of size $h \times w$, denoted as $\mathbf{U} = \{u_{i,j}\}_{i=1,2,\dots,h}^{j=1,2,\dots,w}$, the calculation of DispEn_{2D} involves the following steps:

To map $\mathbf{U}$ into $\mathbf{Y} = \{y_{i,j}\}_{i=1,2,\dots,h}^{j=1,2,\dots,w}$, where $y_{i,j} \in \mathbb{R}$ within the range of $[0, 1]$, selecting a suitable mapping function that accounts for the characteristics of the input data is necessary. Two commonly used mapping approaches include linear mapping and the normal cumulative distribution function (NCDF), often referred to as a sigmoid function. However, given that the distribution of data in natural processes is typically unbalanced, NCDF is frequently preferred.

The NCDF follows the form:

$$y_{i,j}(u_{i,j}) = \frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^{u_{i,j}} \exp \frac{-(t-\mu)^2}{2\sigma^2} dt$$



where μ and σ represent the average and standard deviation of $\mathbf{U}$, respectively. Given that pixel values in image data typically fall within a specific range, the normal cumulative distribution function (NCDF) is chosen as the mapping function for all experiments involving DispEn_{2D} and DispSampEn_{2D}, as detailed in this thesis.

$\mathbf{Y}$ can be linearly mapped into $\mathbf{Z}^c = \{z_{i,j}^c\}_{i=1,2,\dots,h}^{j=1,2,\dots,w}$, where $z_{i,j}^c \in \mathbb{N}$ ranges from 1 to c . This mapping is achieved by applying the following function:

$$z_{i,j}^c(y_{i,j}) = \text{round}(cy_{i,j} + 0.5)$$

Here, c represents the number of classes, a parameter determined by the users, and is often set to 5 in DispEn_{2D} [14].

For each embedding vector $\mathbf{Z}_{i,j}^{c,\mathbf{m}}$ with an embedding dimension vector $\mathbf{m} = [m_h, m_w]$, it is defined as:

$$\mathbf{Z}_{i,j}^{c,\mathbf{m}} = \begin{bmatrix} z_{i,j}^c & \cdots & z_{i,j+(m_w-1)}^c \\ \vdots & \ddots & \vdots \\ z_{i+(m_h-1),j}^c & \cdots & z_{i+(m_h-1),j+(m_w-1)}^c \end{bmatrix}$$

It is asserted that $\mathbf{Z}_{i,j}^{c,\mathbf{m}}$ corresponds to one of the dispersion patterns $\pi_{v_0,\dots,v_{m_h \times m_w - 1}}$. The frequency of occurrence for each pattern can be computed as follows:

$$p(\pi_{v_0,\dots,v_{m_h \times m_w - 1}}) = \frac{[\text{number of } \mathbf{Z}_{i,j}^{c,\mathbf{m}} \text{ where } \mathbf{Z}_{i,j}^{c,\mathbf{m}} = \pi_{v_0,\dots,v_{m_h \times m_w - 1}}]}{(h - (m_h - 1))(w - (m_w - 1))}$$

Finally, the calculation of DispEn_{2D} adopts the form of Shannon entropy:

$$\text{DispEn}_{2D}(\mathbf{U}, \mathbf{m}, c) = - \sum_{\pi=1}^{c^{m_h \times m_w}} p(\pi_{v_0, \dots, v_{m_h \times m_w - 1}}) \ln(p(\pi_{v_0, \dots, v_{m_h \times m_w - 1}}))$$



It is pertinent to note that the total number of potential patterns is $c^{m_h \times m_w}$, given that each element disperses into c classes and the embedding vector size is $m_h \times m_w$.

2.2 Two-dimensional Sample Entropy

In a parallel manner to DispEn_{2D} , SampEn_{2D} denotes the two-dimensional extension of SampEn_{1D} . In the context of an input image $\mathbf{U}$ sized $h \times w$, the definitions of $\phi_{i,j}^{\mathbf{m}}$ and $\phi^{\mathbf{m}}$ are as follows:

$$\phi_{i,j}^{\mathbf{m}}(r) = \frac{[\text{number of } \mathbf{U}_{a,b}^{\mathbf{m}} \text{ where } d[\mathbf{U}_{i,j}^{\mathbf{m}}, \mathbf{U}_{a,b}^{\mathbf{m}}] \leq r \wedge (i,j) \neq (a,b)]}{(h - m_h)(w - m_w) - 1}$$

$$\phi^{\mathbf{m}}(r) = \frac{1}{(h - m_h)(w - m_w)} \sum_{i=1}^{h-m_h} \sum_{j=1}^{w-m_w} \phi_{i,j}^{\mathbf{m}}(r)$$

Here, $\mathbf{U}_{a,b}^{\mathbf{m}}$ denotes the embedding vector, as defined in DispEn_{2D} , utilizing an embedding dimension vector $\mathbf{m} = [m_h, m_w]$. The parameter r , conventionally established at 0.24 times the standard deviation of the input image, functions as the threshold for similarity. $d[\mathbf{U}_{i,j}^{\mathbf{m}}, \mathbf{U}_{a,b}^{\mathbf{m}}]$ represents the maximum absolute value among all elements in the disparity of two embedding vectors [11]. Consequently, $\phi^{\mathbf{m}}$ delineates the ratio of similarity patterns between any two distinct embedding vectors.

The definition of $\phi^{\mathbf{m}+1}$ indeed resembles that of $\phi^{\mathbf{m}}$, albeit with a slight distinction. The total number of comparison pairs $(\mathbf{U}_{i,j}^{\mathbf{m}+1}, \mathbf{U}_{a,b}^{\mathbf{m}+1})$, denoted as the denominator in the

formula for $\phi^{\mathbf{m}+1}$, is $(h - m_h)(w - m_w) \times ((h - m_h)(w - m_w) - 1)$ instead of $(h - m_h - 1)(w - m_w - 1) \times ((h - m_h - 1)(w - m_w - 1) - 1)$. While this adjustment may initially appear perplexing, its purpose is to conveniently nullify the constant terms of $\phi^{\mathbf{m}}$ and $\phi^{\mathbf{m}+1}$ within the fraction.

136	138	164	185	165	189	172
94	86	75	137	120	74	92
71	82	92	71	100	97	21
48	80	131	130	80	91	30
65	85	139	132	69	63	98
141	91	85	71	104	42	77
156	114	76	95	45	0	152

Figure 2.1: Example of first and last embedding vectors for SampEn_{2D}.

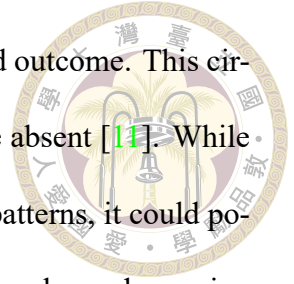
Figure 2.1 depicts the initial (top left) and final (bottom right) embedding vectors employed in a basic input image, with the green frame denoting the computation for $\phi^{\mathbf{m}}$ and the yellow frame for $\phi^{\mathbf{m}+1}$.

Finally, SampEn_{2D} is calculated as follows:

$$SampEn_{2D}(\mathbf{U}, \mathbf{m}, r) = -\ln \frac{\phi^{\mathbf{m}+1}(r)}{\phi^{\mathbf{m}}(r)}$$

SampEn_{2D} is considered a conditional entropy measure as it quantifies the count of similar patterns with the embedding dimension vector $\mathbf{m} + 1$, given that such pairs of embedding vectors are akin with the embedding dimension vector $\mathbf{m}$. It is pertinent to note that SampEn_{2D} entails a computational complexity of $O((h \times w)^2)$ due to the summative processes within $\phi_{i,j}^{\mathbf{m}}$ and $\phi^{\mathbf{m}}$. A notable limitation of the formulation is that either $\phi^{\mathbf{m}}$ or

ϕ^{m+1} must be non-zero; otherwise, SampEn_{2D} may yield an undefined outcome. This circumstance arises in instances where matching embedding vectors are absent [11]. While increasing the threshold r may augment the probability of matching patterns, it could potentially compromise the granularity of information for SampEn_{2D}, thereby underscoring its drawbacks particularly in smaller images or with larger embedding dimensions.





Chapter 3

Two Dimensional

Dispersion Sample Entropy

In this chapter, the proposed algorithm, DispSampEn_{2D} , is introduced, building upon the concept of conditional entropy but offering superior stability and efficiency in comparison to SampEn_{2D} . Furthermore, in contrast to DispEn_{2D} , the method demonstrates heightened robustness against noise signals applied to image textures.

3.1 Methodology

DispSampEn_{2D} maintains the dispersion property of DispEn_{2D} while integrating the concept of conditional entropy akin to SampEn_{2D} . Given an input image denoted as $\mathbf{U} = \{u_{i,j}\}_{i=1,2,\dots,h}^{j=1,2,\dots,w}$, the transformation of $\mathbf{U}$ into $\mathbf{Z}^c = \{z_{i,j}^c\}_{i=1,2,\dots,h}^{j=1,2,\dots,w}$ follows a procedure resembling DispEn_{2D} , where $z_{i,j}^c \in \mathbb{N}$ spans from 1 to c . Here, c represents the number of classes used to disperse the pixel values.

Computed in a manner akin to SampEn_{2D} , portions of DispSampEn_{2D} are defined as

follows:

$$\phi_{i,j}^{\mathbf{m}}(c) = \frac{[\text{number of } \mathbf{Z}_{a,b}^{c,\mathbf{m}} \text{ where } \mathbf{Z}_{i,j}^{c,\mathbf{m}} = \mathbf{Z}_{a,b}^{c,\mathbf{m}} \wedge (i,j) \neq (a,b)]}{(h - m_h)(w - m_w) - 1}$$

$$\phi^{\mathbf{m}}(c) = \frac{1}{(h - m_h)(w - m_w)} \sum_{i=1}^{h-m_h} \sum_{j=1}^{w-m_w} \phi_{i,j}^{\mathbf{m}}(c)$$

and

$$\phi_{i,j}^{\mathbf{m}+1}(c) = \frac{[\text{number of } \mathbf{Z}_{a,b}^{c,\mathbf{m}+1} \text{ where } \mathbf{Z}_{i,j}^{c,\mathbf{m}+1} = \mathbf{Z}_{a,b}^{c,\mathbf{m}+1} \wedge (i,j) \neq (a,b)]}{(h - m_h)(w - m_w) - 1}$$

$$\phi^{\mathbf{m}+1}(c) = \frac{1}{(h - m_h)(w - m_w)} \sum_{i=1}^{h-m_h} \sum_{j=1}^{w-m_w} \phi_{i,j}^{\mathbf{m}+1}(c)$$

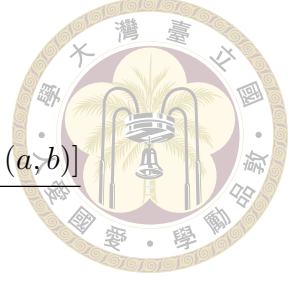
Finally, the calculation of DispSampEn_{2D} proceeds as follows:

$$\text{DispSampEn}_{2D}(\mathbf{U}, \mathbf{m}, c) = -\ln \frac{\phi^{\mathbf{m}+1}(c)}{\phi^{\mathbf{m}}(c)}$$

The crucial aspect here is the utilization of dispersion to replace the range of the similarity threshold in SampEn_{2D} . Consequently, in the definition of counting $\phi_{i,j}^{\mathbf{m}}$, reliance on exact matches between the two embedding vectors is imperative.

3.2 Suggested Value for Dispersion Classes

As outlined in Chapter 2, SampEn_{2D} faces the obstacle of producing undefined values in instances where no similar embedding vector pairs are detected. While widening the threshold range can partially alleviate this issue, selecting an appropriate threshold poses a significant challenge. Consequently, the prevalent approach is to utilize 0.24 times the standard deviation. In contrast, DispSampEn_{2D} employs dispersion classes to supplant the threshold. The advantage of this approach lies in the ability to circumvent this issue under



specific conditions by selecting an appropriate parameter c .

Recalling that the total number of potential patterns with the embedding dimension vector $(\mathbf{m} + 1)$ after the dispersion process is $c^{(m_h+1) \times (m_w+1)}$, and the total number of embedding vectors is $(h - m_h)(w - m_w)$, the Pigeonhole Principle can be applied. According to this principle, there must be at least one pair of matching embedding vectors when the following condition is satisfied:

$$(h - m_h) \times (w - m_w) > c^{(m_h+1) \times (m_w+1)}$$

Based on the inequality provided above, a suggested value for the dispersion classes c is proposed as follows:

$$c_{\text{suggest}} = \max \left(2, \left\lfloor \sqrt{(m_h+1) \times (m_w+1)} \right\rfloor \right)$$

3.3 Linear Approach for Computation

As previously discussed, the naive approach of SampEn_{2D} becomes an $O((m_h \times m_w)^2)$ algorithm due to the recursive loop for counting $\phi^{\mathbf{m}}$. This can result in significant computational burden, particularly with large-size images. However, despite the necessity of computing $\phi^{\mathbf{m}}$ in DispSampEn_{2D} , an implementation is proposed that operates in $O(m_h \times m_w)$ time complexity.

Recall that

$$\phi^{\mathbf{m}}(c) = \frac{1}{N} \sum_{i=1}^{h-m_h} \sum_{j=1}^{w-m_w} [\text{number of } \mathbf{Z}_{a,b}^{c,\mathbf{m}} \text{ where } \mathbf{Z}_{i,j}^{c,\mathbf{m}} = \mathbf{Z}_{a,b}^{c,\mathbf{m}} \wedge (i,j) \neq (a,b)]$$

where N is a constant equal to $((h - m_h)(w - m_w)) \times ((h - m_h)(w - m_w) - 1)$. Since DispSampEn_{2D} counts the number of similarity embedding vectors by exact matching of two different embedding vectors, the computation can be simplified by counting the number of embedding vectors that are equal to specific patterns, similar to what is done in DispEn_{2D} .

Hence, the following formula holds:

$$\text{DispSampEn}_{2D}(\mathbf{U}, \mathbf{m}, c) = \frac{1}{N} \sum_{\pi=1}^{c^{m_h \times m_w}} \text{Count}_{\pi} \times (\text{Count}_{\pi} - 1)$$

where

$$\text{Count}_{\pi} = [\text{number of } \mathbf{Z}_{i,j}^{c,\mathbf{m}} \text{ where } \mathbf{Z}_{i,j}^{c,\mathbf{m}} = \pi_{v_0, \dots, v_{m_h \times m_w - 1}}]$$

The equations above indicate that for every embedding vector equal to the pattern π , there exist $(\text{Count}_{\pi} - 1)$ other similar embedding vectors. Therefore, we can compute DispSampEn_{2D} in linear time, a property crucial for its practical applicability.



Chapter 4 Datasets

In this chapter, the synthetic and real-world datasets employed in our experiments to validate and evaluate our proposed metric, DispSampEn_{2D}, are introduced. It is crucial to highlight that all images utilized in this study are of size 128x128 pixels.

4.1 Synthetic Datasets

For validation purposes, several images with varying levels of irregularity were synthesized using the two-dimensional MIX process (MIX_{2D}(p)). The sinusoidal function was chosen as the deterministic component, while uniform white noise served as the stochastic component.

MIX_{2D}(p) extends the one-dimensional MIX process [18], incorporating both a deterministic signal and a stochastic signal. It is defined as follows:

$$\text{MIX}_{2D}(p)_{i,j} = (1 - Z_{i,j})X_{i,j} + Z_{i,j}Y_{i,j}$$

Here, the parameter $p \in [0, 1]$ governs the degree of irregularity in the resultant images, with higher values of p corresponding to greater irregularity. $Z_{i,j}$ is a random variable that assumes a value of 1 with probability p and 0 with probability $(1 - p)$. $X_{i,j}$ represents the

deterministic component, defined as $X_{i,j} = \sin(\frac{2\pi i}{12}) + \sin(\frac{2\pi j}{12})$ in our case. Conversely, $Y_{i,j}$ represents the stochastic component, indicating uniform white noise within the range $[-\sqrt{3}, \sqrt{3}]$ that we applied [11]. The results are depicted in Figure 4.1, illustrating the increase in irregularity with the increment of p as expected.

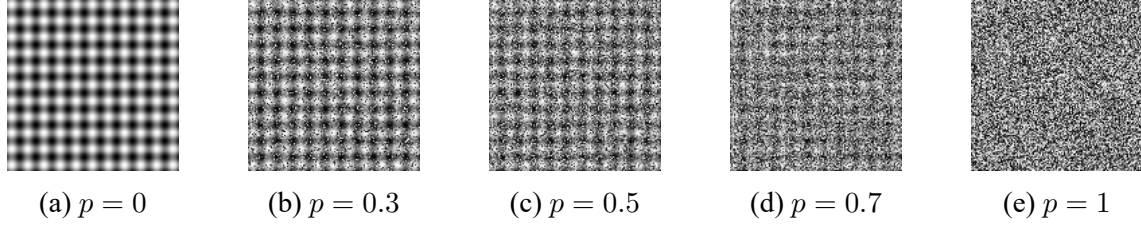


Figure 4.1: Results of MIX_{2D} with different p .

4.2 Real-world Datasets

Two real-world datasets are utilized to evaluate the performance of our metric in classification tasks in Chapter 5. The first dataset employed is the normalized Brodatz texture dataset. This dataset comprises 112 grayscale real-world images. Consistent with prior studies, 9 groups of images are selected for our experiments [11, 14]. For each group, 81 samples are extracted. Sample images from this dataset are illustrated in Figure 4.2.

The second dataset is the Kylberg texture dataset. While the Brodatz dataset has been widely used in texture analysis studies, the Kylberg dataset provides more recent and higher-quality images. In the Kylberg dataset, 10 groups of textures are selected, with each group containing 160 samples [14]. These samples are depicted in Figure 4.3.

Additionally, MIX_{2D} is applied to these two datasets for the noise resistance test in Chapter 5. Figure 4.4 and 4.5 depict the D5 group and blanket1 group, respectively, with varying levels of noise mixed in.

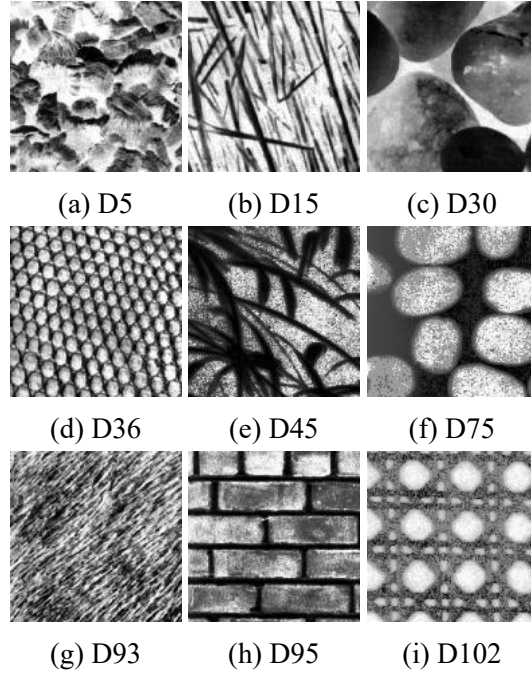


Figure 4.2: Samples from each group used in the normalized Brodatz dataset.

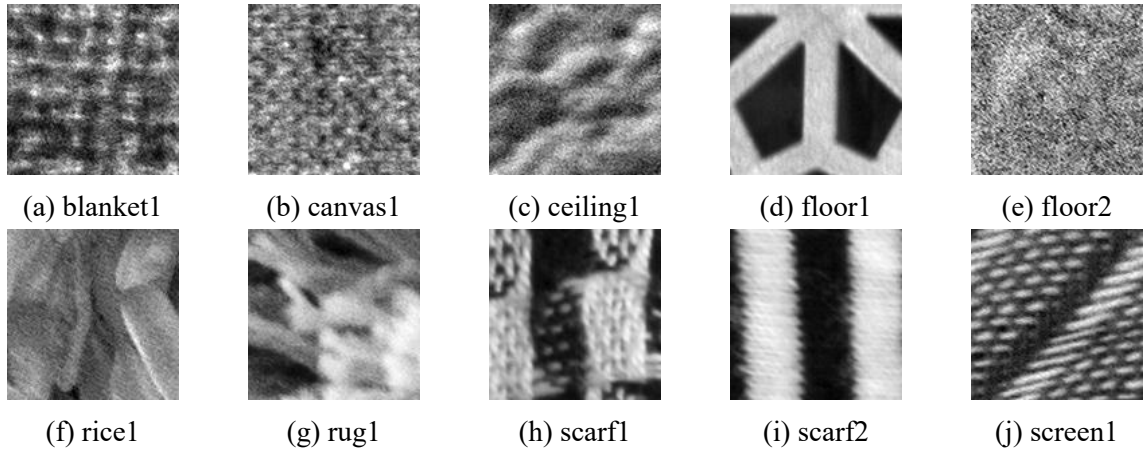


Figure 4.3: Samples from each group used in the Kylberg dataset.

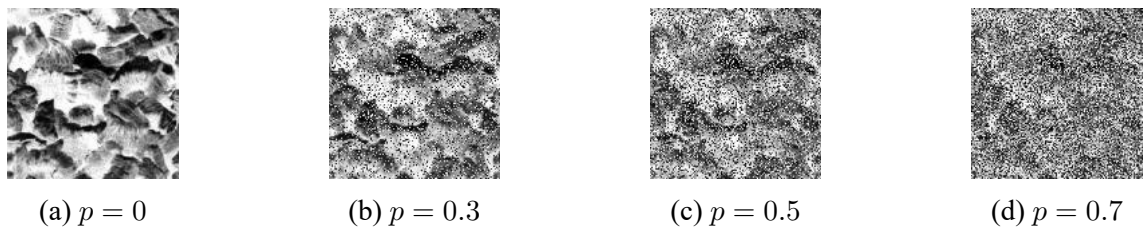


Figure 4.4: MIX_{2D} applied on D5 with varying p .

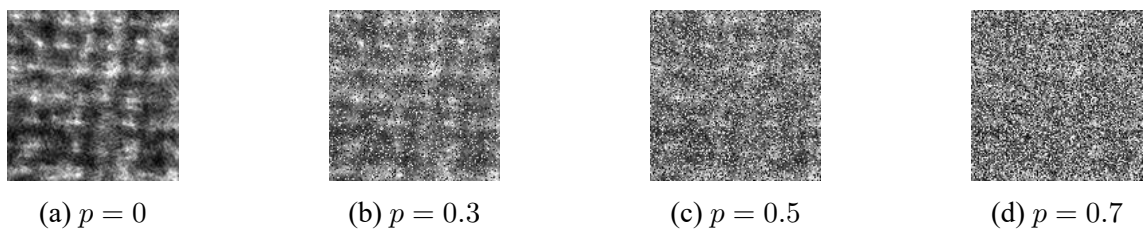


Figure 4.5: MIX_{2D} applied on blanket1 with varying p .



Chapter 5 Experiments and Results

In this chapter, a series of experiments are designed to evaluate the validation and usability of DispSampEn_{2D} . It is important to note that a value of 0.24 times the standard deviation is utilized as the threshold of similarity r for SampEn_{2D} [11], the number of classes c is set to 5 for DispEn_{2D} [14], and the suggested value of dispersion classes as delineated in Chapter 3 is adopted for DispSampEn_{2D} . Additionally, square embedding vectors are applied; therefore, for the instance when $m = 2$, it actually implies $\mathbf{m} = [2, 2]$ in this chapter.

5.1 Validation of DispSampEn_{2D}

In this experiment, the objective is to assess the capability of DispSampEn_{2D} to discern the degree of irregularity within a provided image, and to contrast the trends observed in the line charts as irregularity increases with DispEn_{2D} and SampEn_{2D} .

Accordingly, the synthesis dataset and the MIX process with varying p from 0 to 1, utilizing an interval of 0.1, are employed. The experiment outputs entropy values for different metrics, with the results illustrated in Figure 5.1.

Observing the line charts, it becomes apparent that entropy for all metrics generally

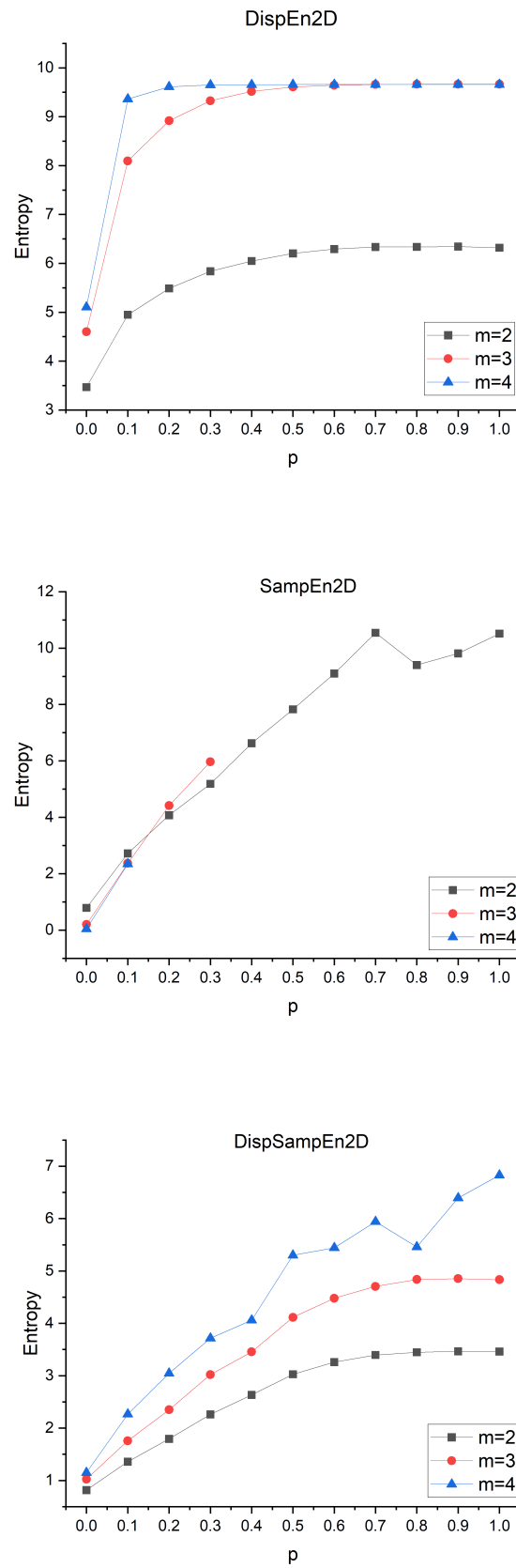


Figure 5.1: Entropy of different metrics plotted against the value p of the MIX process.

increases as p increases. Furthermore, there are additional observations. Firstly, DispEn_{2D} exhibits a sharp rise with lower p but experiences minimal change with higher p , as it reaches saturation entropy value prematurely. Consequently, this phenomenon renders DispEn_{2D} less effective in distinguishing images with varying levels of irregularity under high-noise conditions.

Secondly, SampEn_{2D} encounters the issue of yielding undefined values as p increases. Moreover, this issue exacerbates when the embedding dimension m becomes larger. This represents the inherent drawback discussed in Chapter 2.

Finally, DispSampEn_{2D} demonstrates a strong correlation with irregularity without encountering undefined values. Notably, it circumvents the drawbacks of both DispEn_{2D} and SampEn_{2D} , rendering it a promising choice for texture analysis. These enhancements align with our initial expectations when designing the metrics.

5.2 Performance on Brodatz dataset

In this experiment, DispSampEn_{2D} is applied to the Brodatz dataset as mentioned in Chapter 4, aiming to validate its performance on real-world data. Furthermore, to assess the usability of DispSampEn_{2D} , a naive Bayes classifier is employed, utilizing a single entropy value as input to evaluate the accuracy in classifying different groups of images. Additionally, the performance is compared with that of DispEn_{2D} and SampEn_{2D} .

DispSampEn_{2D} is computed for each group and compared with the other two entropy metrics, the results are depicted in Figure 5.2. It is noteworthy that each group comprises 81 samples in our Brodatz dataset. Consequently, the mean is annotated along with the values added or subtracted by one standard deviation. Notably, we observe that

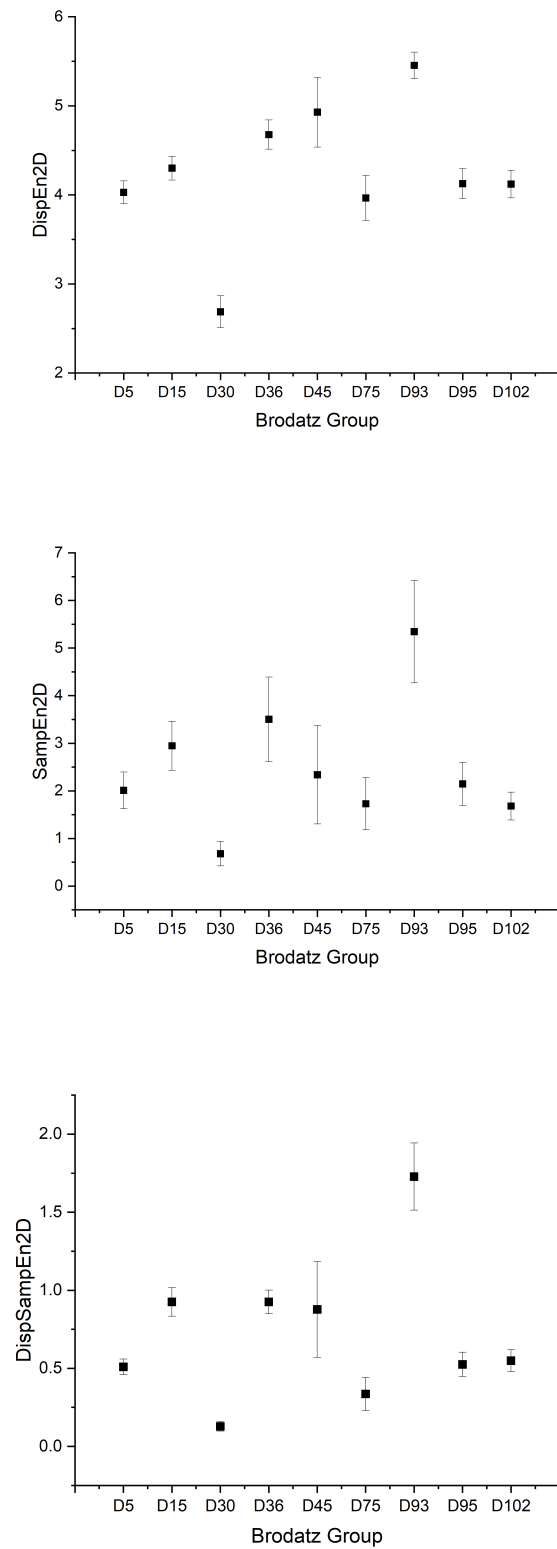
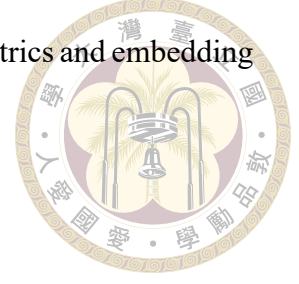


Figure 5.2: Entropy for each group in the Brodatz dataset.

Table 5.1: Accuracy of classifying the Brodatz dataset for different metrics and embedding dimensions.

	m=2	m=3	m=4
DispEn _{2D}	53.42%	51.60%	53.42%
SampEn _{2D}	46.57%	41.55%	34.25%
DispSampEn _{2D}	56.62%	57.99%	58.90%



D45 and D93 exhibit higher standard deviation compared to others, while D15 and D36 demonstrate similar entropy values, as do D95 and D102. Additionally, DispSampEn_{2D} and DispEn_{2D} show less overlap among the groups compared to SampEn_{2D}, implying the better performance on the Brodatz dataset.

Finally, the accuracy of the classification task for different entropy metrics is compared, and the results are presented in Table 5.1. It is noteworthy that, on the Brodatz dataset, DispSampEn_{2D} exhibits the best performance across all different embedding dimensions, followed by DispEn_{2D} in second place, and SampEn_{2D} in third.

5.3 Performance on Kylberg dataset

In this experiment, similar procedures are conducted as in the previous one, with the sole distinction being the utilization of the Kylberg dataset instead of the Brodatz dataset. Consequently, the results of applying DispSampEn_{2D} on the Kylberg dataset are depicted in Figure 5.3. Notably, it is observed that the Kylberg dataset exhibits lower standard deviation for each group compared to the Brodatz dataset. Additionally, the overlap among groups is considerably reduced in the Kylberg dataset compared to Brodatz, which is advantageous for the classification task.

The accuracy of each metric is also presented in Table 5.2. As anticipated, all three metrics exhibited superior performance on the Kylberg dataset compared to the Brodatz

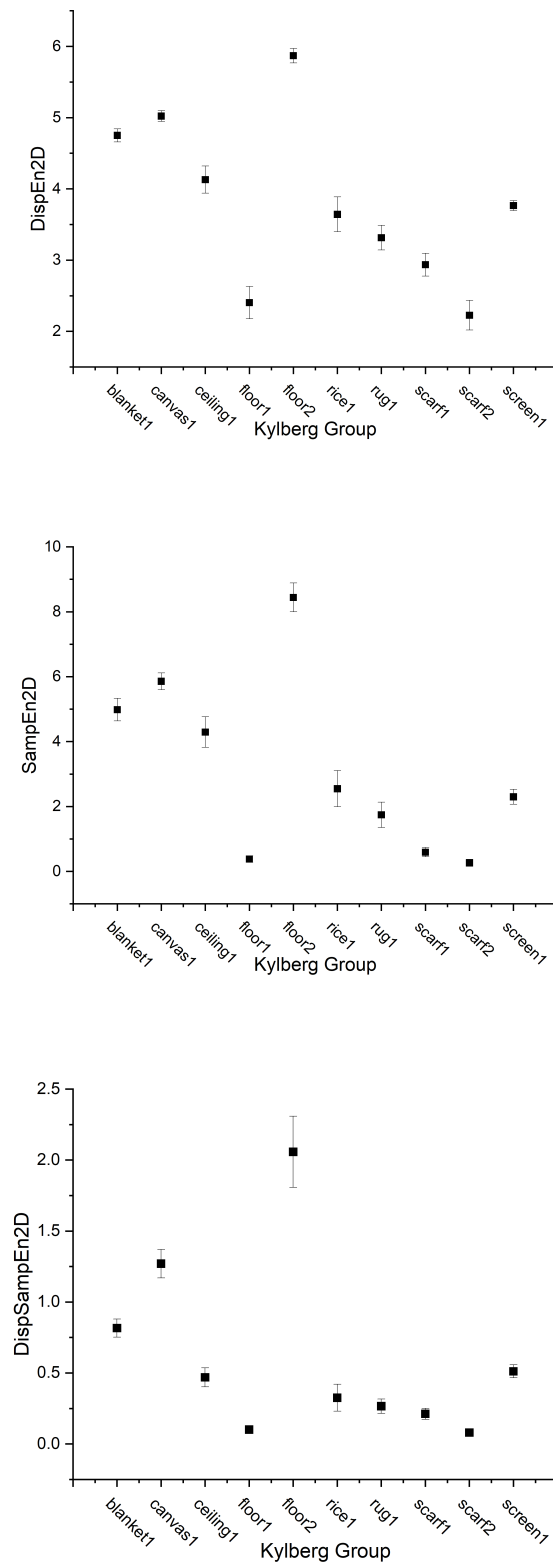
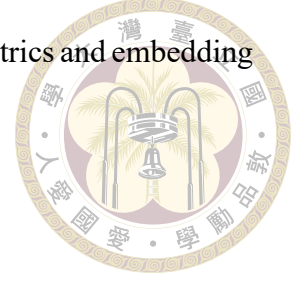


Figure 5.3: Entropy for each group in the Kylberg dataset.

Table 5.2: Accuracy of classifying the Kylberg dataset for different metrics and embedding dimensions.

	m=2	m=3	m=4
DispEn _{2D}	78.33%	78.54%	69.58%
SampEn _{2D}	73.96%	57.92%	51.88%
DispSampEn _{2D}	75.21%	73.75%	68.75%



dataset. However, DispEn_{2D} demonstrated the highest performance on the Kylberg dataset, closely followed by DispSampEn_{2D}, while SampEn_{2D} lagged significantly behind.

5.4 Noise Resistance

In this experiment, both the Brodatz and Kylberg datasets are subjected to varying degrees of noise added by the MIX process. Additionally, the naive Bayes classifier is employed to assess the noise resistance capabilities of each entropy metric. Several intriguing findings emerged from this experiment, and further insights are proposed in the subsequent section.

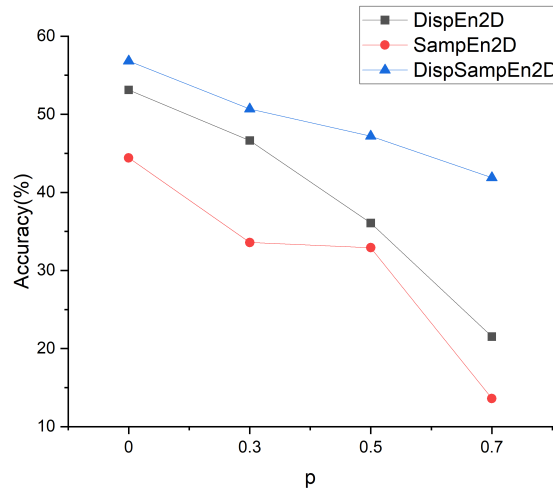


Figure 5.4: Accuracy of classifying the Brodatz dataset with varying p for different metrics, with $m = 2$.

Figure 5.4 and Figure 5.5 display the classifying accuracy for DispEn_{2D}, SampEn_{2D},

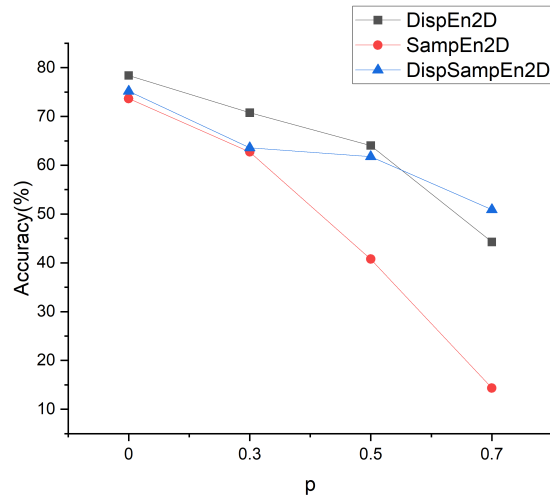


Figure 5.5: Accuracy of classifying the Kylberg dataset with varying p for different metrics, with $m = 2$.

and DispSampEn_{2D} on the Brodatz and Kylberg datasets, respectively. The noise is introduced by the MIX process with p varying across 0, 0.3, 0.5, and 0.7. Notably, as p increases, the accuracy of SampEn_{2D} decreases rapidly, while DispEn_{2D} performs slightly better. Remarkably, DispSampEn_{2D} maintains significantly higher accuracy compared to the other two metrics, especially under high-noise conditions.

To investigate the reasons behind the inferior performance of DispEn_{2D} and SampEn_{2D}, further experiments are conducted. Firstly, DispEn_{2D} for each group from both the Brodatz and Kylberg datasets with p equaling 0 and 0.7 are computed, respectively. The results are presented in Figure 5.6 and Figure 5.7. It is evident that, for both the Brodatz and Kylberg datasets, when significant noise is applied, the overlaid range of DispEn_{2D} for different groups becomes wider. This phenomenon directly results in the difficulty of distinguishing images from different groups. The underlying cause is trivial; it is the tendency of DispEn_{2D} to reach saturation entropy value prematurely, as found in Section 5.1.

As for SampEn_{2D}, the ratios of the occurrences of undefined values for the Brodatz

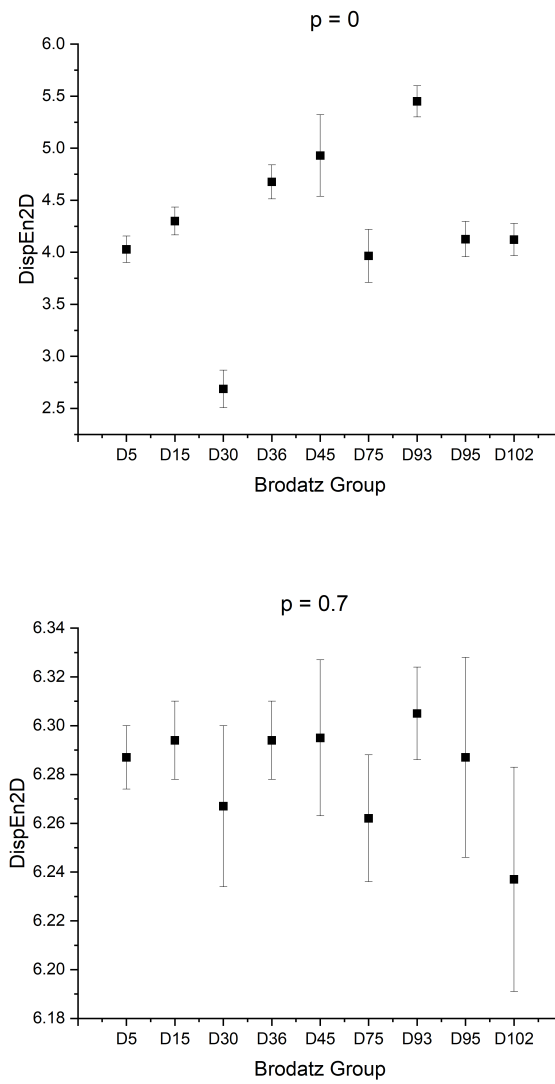


Figure 5.6: DispEn_{2D} for each group from the Brodatz dataset.

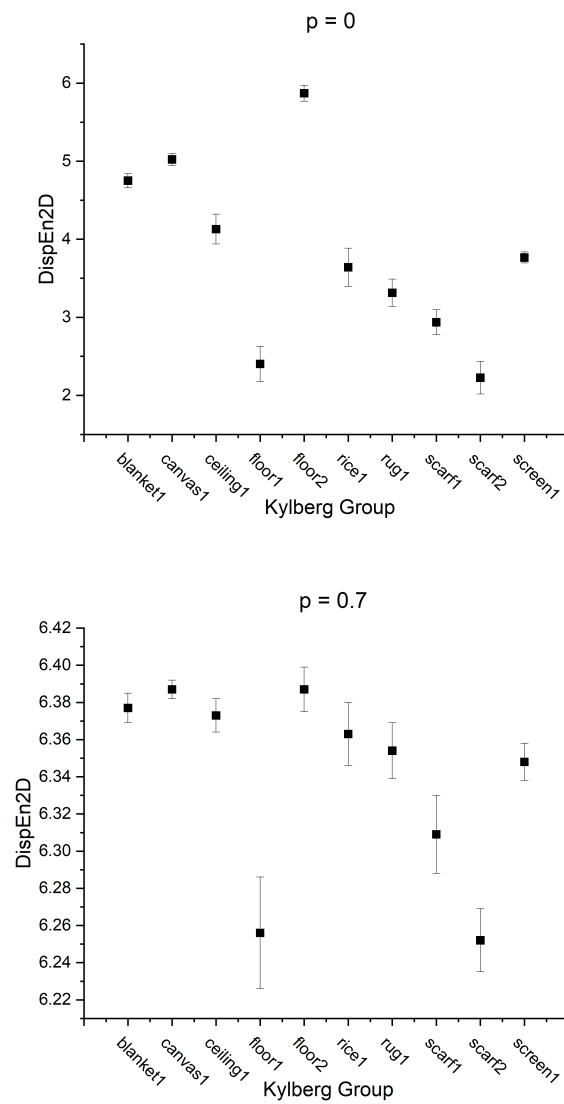


Figure 5.7: DispEn_{2D} for each group from the Kylberg dataset.

Table 5.3: Accuracy and respective undefined value ratios for SampEn_{2D} on the noisy Brodatz datasets.

	p=0	p=0.3	p=0.5	p=0.7
Undefined value ratio	0%	0%	0%	8.23%
Accuracy	44.42%	33.59%	32.92%	13.61%

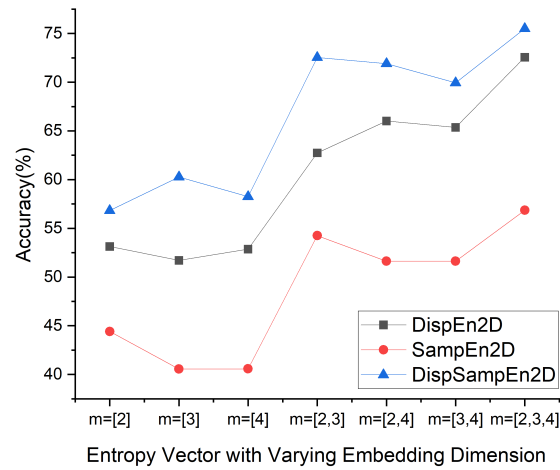
Table 5.4: Accuracy and respective undefined value ratios for SampEn_{2D} on the noisy Kylberg datasets.

	p=0	p=0.3	p=0.5	p=0.7
Undefined value ratio	0.06%	0.19%	1.06%	13.31%
Accuracy	73.67%	67.72%	40.76%	14.39%

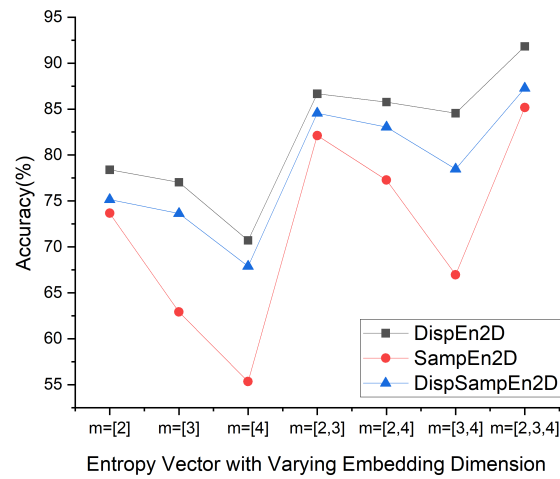
and Kylberg datasets with different degrees of noise are calculated. The correlation between the accuracy and the undefined value ratios is presented in Table 5.3 and Table 5.4. It is noticeable that the accuracy drops dramatically while p ranges from 0.5 to 0.7 on the Brodatz dataset and from 0.3 to 0.7 on the Kylberg dataset, with undefined value ratios increasing rapidly within these ranges. Additionally, all undefined values are set to zero during the data preprocessing before classification in our experiment. Hence, it is apparent that the issue of undefined values, inherent to SampEn_{2D}, detrimentally impacts the classification task performance.

5.5 Entropy Vector

The embedding dimension is a crucial parameter for the entropy metrics and is extensively discussed in the majority of related studies. However, there is no universal method to determine the embedding dimension for all purposes; the optimal embedding dimension is determined on a case-by-case basis and often through trial and error. For the classification experiments in previous related studies and the sections above, a single entropy value with the determined dimension is utilized as the input of the classifier [11, 12, 14].

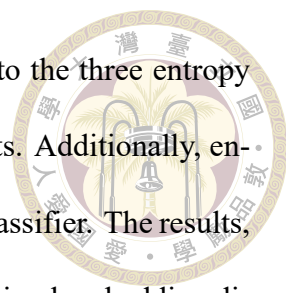


(a) Brodatz



(b) Kylberg

Figure 5.8: Accuracy on real-world datasets, using entropy vectors with varying embedding dimensions as classifier input.



In this experiment, various embedding dimensions are applied to the three entropy metrics to evaluate accuracy on both the Brodatz and Kylberg datasets. Additionally, entropy vectors with different parameters are proposed as input for the classifier. The results, depicted in Figure 5.8, reveal several noteworthy findings. Firstly, optimal embedding dimension choices differ for each metric. For example, $m = 2$ is optimal for DispEn_{2D} and SampEn_{2D} on the Brodatz dataset, while $m = 3$ is preferable for DispSampEn_{2D}. Secondly, the optimal parameters also vary between the datasets; notably, DispSampEn_{2D} with $m = 2$ performs best on the Kylberg dataset, contrary to the $m = 3$ preference on the Brodatz dataset.

The aforementioned observations align with the previous study. However, this section reveals a more surprising finding: the results indicate that vectors containing more entropy with varying embedding dimensions as input provide the classifier with better classification capabilities than applying a single entropy value with a predetermined embedding dimension. To uncover the underlying reason for this phenomenon, confusion matrices for using entropy vectors with $m = [2]$, $m = [3]$, $m = [4]$, and $m = [2, 3, 4]$ are plotted in Figure 5.9 and Figure 5.10. When observing the D45 row in Figure 5.9 and the screen1 row in Figure 5.10, it is apparent that $m = [2]$, $m = [3]$, and $m = [4]$ make some mistakes in distinguishing and mispredict them as different groups. However, $m = [2, 3, 4]$ takes advantage among the three to achieve better correct predictions. Accuracy per group for using different entropy vectors on the Brodatz and Kylberg datasets, which corresponds to the diagonal values in the confusion matrices, is depicted in Figure 5.11. Focusing on the green lines $m = [2, 3, 4]$ in the plots, it is observed that accuracy per group is either in the first or second place. It is hypothesized that this is caused by the additive effect of entropy with different embedding dimensions, as revealed in our findings

Table 5.5: Computation time required for calculating the Brodatz dataset for different metrics.

	DispEn _{2D}	SampEn _{2D}	DispSampEn _{2D}
Computation Time(Sec)	5.76	2013.94	8.07

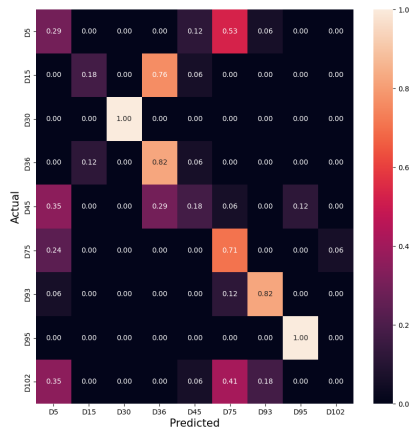
Table 5.6: Computation time required for calculating the Kylberg dataset for different metrics.

	DispEn _{2D}	SampEn _{2D}	DispSampEn _{2D}
Computation Time(Sec)	11.42	5584.75	13.54

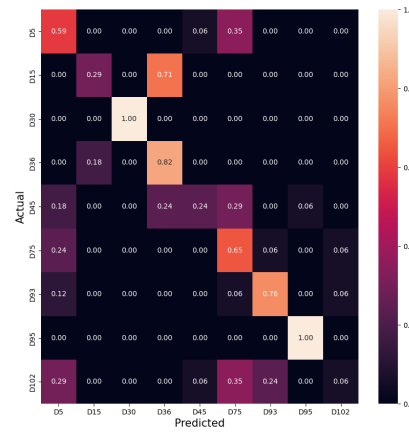
from the confusion matrices.

5.6 Computation Time

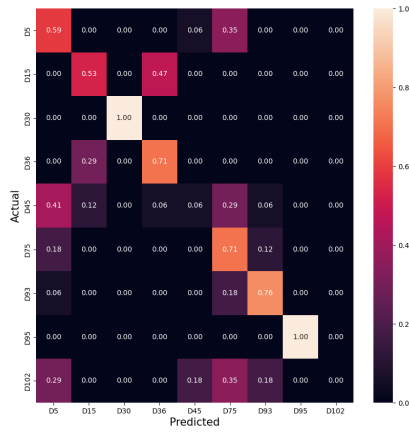
In this section, the computation time is discussed. Referring to the theoretical analysis in the previous chapter, both DispEn_{2D} and DispSampEn_{2D} are implemented in linear time, whereas SampEn_{2D} is a square-time algorithm. The computation time for calculating each entropy metric on the Brodatz and Kylberg datasets is presented in Table 5.5 and Table 5.6, respectively. The calculation process was executed on a single core of AMD Ryzen5 3500X. The results validate our analysis, showing that DispEn_{2D} and DispSampEn_{2D} exhibit significantly better efficiency than SampEn_{2D}, rendering them promising metrics for real-world applications in texture assessment.



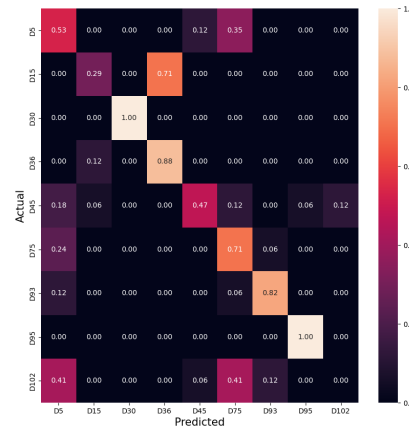
(a) $m = [2]$



(b) $m = [3]$

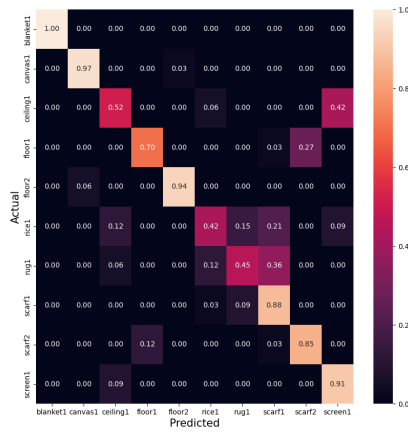


(c) $m = [4]$

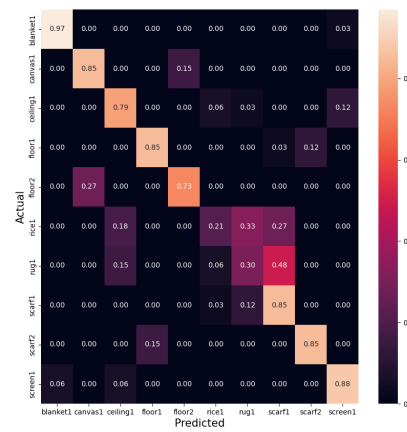


(d) $m = [2, 3, 4]$

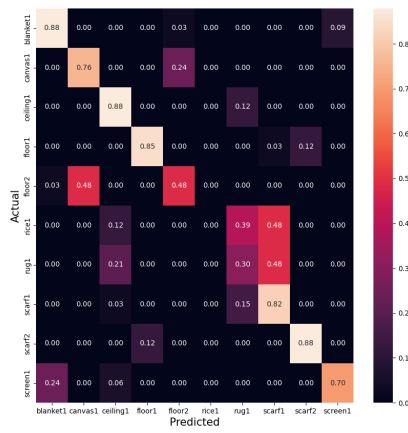
Figure 5.9: Confusion matrices for using different entropy vectors on the Brodatz dataset.



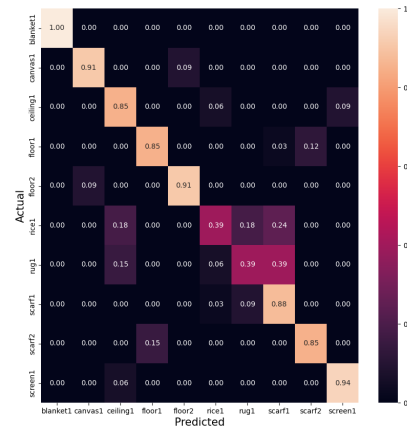
(a) $m = [2]$



(b) $m = [3]$

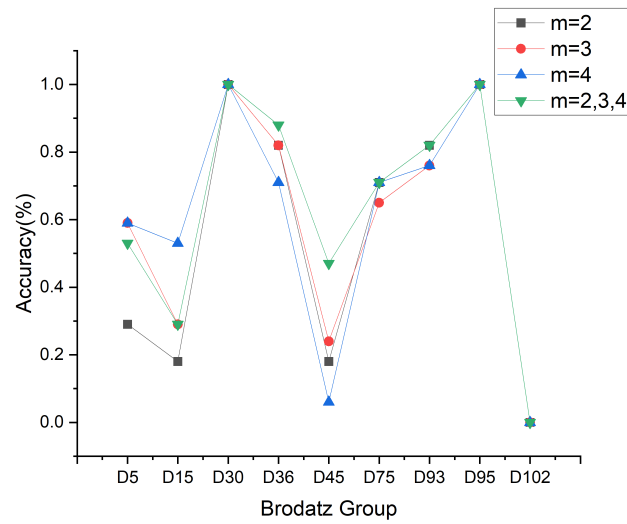


(c) $m = [4]$

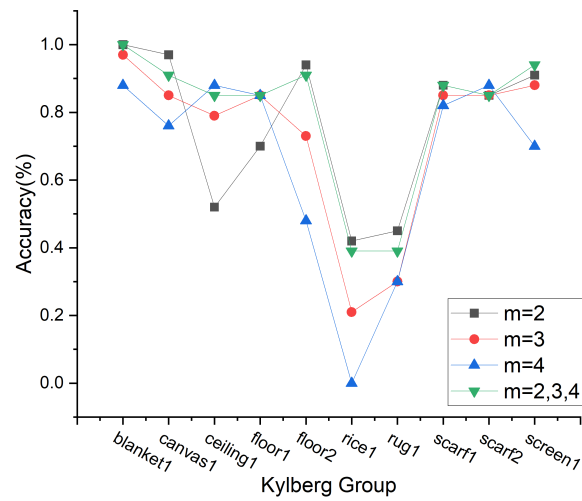


(d) $m = [2, 3, 4]$

Figure 5.10: Confusion matrices for using different entropy vectors on the Kylberg dataset.



(a) Brodatz



(b) Kylberg

Figure 5.11: Accuracy per group for using different entropy vectors on the Brodatz and Kylberg datasets.



Chapter 6 Conclusion

The study introduces a new entropy-based metric called two-dimensional dispersion sample entropy (DispSampEn_{2D}) for evaluating image textures. Derived from conditional entropy, DispSampEn_{2D} ensures avoidance of the undefined value problem encountered in SampEn_{2D} under specific conditions, offering a straightforward linear implementation. Moreover, it exhibits remarkable robustness against noise.

As entropy-based metrics gain traction in texture analysis, numerous metrics have emerged in previous research. To highlight the advantages of the new metric, several experiments were conducted, comparing it with two popular metrics, DispEn_{2D} and SampEn_{2D} . Synthetic datasets and two real-world datasets were utilized to validate and evaluate DispSampEn_{2D} , taking into account its noise resistance and computational efficiency. In classification tasks on real-world datasets, DispSampEn_{2D} performed similarly to DispEn_{2D} and surpassed SampEn_{2D} . Additionally, DispSampEn_{2D} exhibited superior noise resistance, and the computational efficiency of DispSampEn_{2D} is consistent with theoretical expectations.

In addition to the new metric discussed above, the concept of an entropy vector is introduced for application in classifiers. It has been demonstrated that utilizing an entropy vector with various embedding dimensions can enhance accuracy in classification tasks on real-world datasets. This improvement is attributed to the additive effect of different

embedding dimensions.

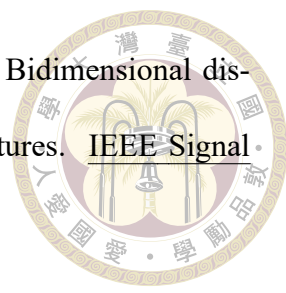
For future research, extending DispSampEn_{2D} to its one-dimensional version and validating it on time-series data would be beneficial. Additionally, exploring whether the additive effect among different types of entropy metrics exists in entropy vectors warrants further investigation.



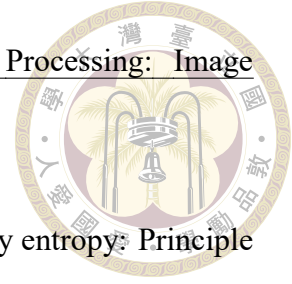


References

- [1] Peng Li, Chengyu Liu, Ke Li, Dingchang Zheng, Changchun Liu, and Yinglong Hou. Assessing the complexity of short-term heartbeat interval series by distribution entropy. Medical & biological engineering & computing, 53:77–87, 2015.
- [2] Christoph Bandt and Bernd Pompe. Permutation entropy: a natural complexity measure for time series. Physical review letters, 88(17):174102, 2002.
- [3] Mostafa Rostaghi and Hamed Azami. Dispersion entropy: A measure for time-series analysis. IEEE Signal Processing Letters, 23(5):610–614, 2016.
- [4] Steven M Pincus. Approximate entropy as a measure of system complexity. Proceedings of the national academy of sciences, 88(6):2297–2301, 1991.
- [5] Joshua S Richman, Douglas E Lake, and J Randall Moorman. Sample entropy. In Methods in enzymology, volume 384, pages 172–184. Elsevier, 2004.
- [6] Alfonso Delgado-Bonal and Alexander Marshak. Approximate entropy and sample entropy: A comprehensive tutorial. Entropy, 21(6):541, 2019.
- [7] Tao Zhang, Wanzhong Chen, and Mingyang Li. Fuzzy distribution entropy and its application in automated seizure detection technique. Biomedical Signal Processing and Control, 39:360–377, 2018.

- 
- [8] Hamed Azami, Javier Escudero, and Anne Humeau-Heurtier. Bidimensional distribution entropy to analyze the irregularity of small-sized textures. IEEE Signal Processing Letters, 24(9):1338–1342, 2017.
- [9] Andreia S Gaudêncio, Mirvana Hilal, João M Cardoso, Anne Humeau-Heurtier, and Pedro G Vaz. Texture analysis using two-dimensional permutation entropy and amplitude-aware permutation entropy. Pattern Recognition Letters, 159:150–156, 2022.
- [10] Luiz Eduardo Virgilio da Silva, Antonio Carlos Da Silva Senra Filho, Valéria Paula Sassoli Fazan, Joaquim Cezar Felipe, and Luiz Otavio Murta. Two-dimensional sample entropy analysis of rat sural nerve aging. In 2014 36th Annual International Conference of the IEEE Engineering in Medicine and Biology Society, pages 3345–3348. IEEE, 2014.
- [11] Luiz Eduardo Virgili Silva, ACS Senra Filho, Valéria Paula Sassoli Fazan, Joaquim Cezar Felipe, and LO Murta Junior. Two-dimensional sample entropy: Assessing image texture through irregularity. Biomedical Physics & Engineering Express, 2(4):045002, 2016.
- [12] Luiz EV Silva, Juliano J Duque, Joaquim C Felipe, Luiz O Murta Jr, and Anne Humeau-Heurtier. Two-dimensional multiscale entropy analysis: Applications to image texture evaluation. Signal Processing, 147:224–232, 2018.
- [13] Ricardo Espinosa, Raquel Bailón, and Pablo Laguna. Two-dimensional espen: A new approach to analyze image texture by irregularity. Entropy, 23(10):1261, 2021.
- [14] Hamed Azami, Luiz Eduardo Virgilio da Silva, Ana Carolina Mieko Omoto, and Anne Humeau-Heurtier. Two-dimensional dispersion entropy: An information-

theoretic method for irregularity analysis of images. Signal Processing: Image Communication, 75:178–187, 2019.



- [15] Mirvana Hilal and Anne Humeau-Heurtier. Bidimensional fuzzy entropy: Principle analysis and biomedical applications. In 2019 41st Annual International Conference of the IEEE Engineering in Medicine and Biology Society (EMBC), pages 4811–4814. IEEE, 2019.
- [16] Cristina Morel and Anne Humeau-Heurtier. Multiscale permutation entropy for two-dimensional patterns. Pattern Recognition Letters, 150:139–146, 2021.
- [17] Mirvana Hilal, Clémence Berthin, Ludovic Martin, Hamed Azami, and Anne Humeau-Heurtier. Bidimensional multiscale fuzzy entropy and its application to pseudoxanthoma elasticum. IEEE Transactions on Biomedical Engineering, 67(7):2015–2022, 2019.
- [18] Steven M Pincus and Ary L Goldberger. Physiological time-series analysis: what does regularity quantify? American Journal of Physiology-Heart and Circulatory Physiology, 266(4):H1643–H1656, 1994.