



Mathematical morphology tools to evaluate periodic linguistic summaries

Gilles Moyse, Marie-Jeanne Lesot, Bernadette Bouchon-Meunier

► To cite this version:

Gilles Moyse, Marie-Jeanne Lesot, Bernadette Bouchon-Meunier. Mathematical morphology tools to evaluate periodic linguistic summaries. Flexible Query Answering Systems, Sep 2013, Granada, Spain. pp.257-268, 10.1007/978-3-642-40769-7_23 . hal-00932844

HAL Id: hal-00932844

<https://hal.science/hal-00932844>

Submitted on 5 Jan 2015

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Mathematical morphology tools to evaluate periodic linguistic summaries

Gilles Moyse, Marie-Jeanne Lesot, Bernadette Bouchon-Meunier

UPMC Univ. Paris 06, CNRS UMR 7606, LIP6, F-75005, Paris, France
`surname.name@lip6.fr`

Abstract This paper considers the task of establishing periodic linguistic summaries of the form “Regularly, the data take high values”, enriched with an estimation of the period and a linguistic formulation. Within the framework of methods that address this task testing whether the dataset contains regularly spaced groups of high and low values with approximately constant size, it proposes a mathematical morphology (MM) approach based on watershed. It compares the proposed approach to other MM methods in an experimental study based on artificial data with different forms and noise types.

Keywords: Fuzzy linguistic summaries, Periodicity computing, Mathematical Morphology, Temporal data mining, Watershed

1 Introduction

Linguistic summaries aim at building human understandable representations of datasets, thanks to natural language sentences. They take different forms representing different kinds of patterns [27,28,12]. In this paper we consider this task in the case of time series for which regularity is looked for, more precisely summaries of the form “Regularly, the data take high values”. If the data are membership degrees to a fuzzy modality A , the sentence can be interpreted as “regularly, the data are A ”. Moreover, if the sentence holds, a candidate period is computed and an appropriate linguistic formulation is generated, based on the choice of a relevant time unit, approximation and adverb. The final sentence can for instance be “Approximately every 20 hours, the data take high values”.

The Detection of Periodic Events (DPE) methodology [21] defines a framework to address this task relying on the assumption that if a dataset contains regularly spaced high and low value groups of approximately constant size, then it is periodic. It consists in 3 steps: clustering, cluster size regularity and linguistic rendering. In [21], the first step is based on the calculation of an erosion score based on Mathematical Morphology [24]. In this paper, we propose to apply the DPE methodology using a new clustering method depending on a watershed approach [4] and to compare it in an enriched experimental protocol.

Section 2 presents an overview of related works. A reminder about the evaluation of periodic protoforms is given in Section 3. The proposed watershed method is described in Section 4. Lastly, Section 5 presents experimental results on artificial data comparing the two approaches as well as a baseline method.

2 Related Works

This section briefly describes the principles of linguistic summaries, temporal data mining and period detection in signal processing, at the crossroads of which the considered DPE methodology lies. To the best of our knowledge, DPE is the first approach combining these fields.

2.1 Linguistic Summaries

Linguistic summaries aim at building compact representations of datasets, in the form of natural language sentences describing their main characteristics. Besides approaches based on natural language generation techniques, they can be produced using fuzzy logic, in which case they are called fuzzy linguistic summaries (see [5,13] for a comparison between these two areas).

Introduced in the seminal papers [11,27,28], they are built on sentences called “protoforms”, such as “ QX are A ” where Q is a quantifier (e.g. “most” or “around 10”), A a linguistic modality associated with one of the attributes (e.g. “young” for the attribute “age”) and X the data to summarise. The relevance of a candidate protoform, measured by the truth degree of its instantiation for the considered data, depends on the Σ -count of the dataset according to the chosen fuzzy modality. Extensions have been defined to handle the temporal nature of data, using a “Trend” attribute [10] or considering fuzzy temporal propositions [6] to restrict the truth value of a summary to a certain period of time, but they do not cope with periodicity.

2.2 Temporal Data Mining

Temporal data mining is a domain that groups various issues related to data mining taking into account the temporal aspect of the data (see [8,15] for exhaustive states of the art). Some methods aim at discovering frequent patterns, using extensions of the Apriori algorithm [1], possibly dedicated to long sequences [19] or with time-window or duration constraints [16,20]. Although mining recurring events, these approaches are not concerned with periodicity. Cyclic association rules [22] are satisfied on a fixed periodic basis: the time axis is split into constant length segments against which association rules are tested. So as to automatically compute a candidate period, extensions based on a Fourier transform [3] or a statistical test over the average interval between events [17] can be used.

2.3 Signal Processing for Period Detection

Period detection is a well known problem in signal processing and several methods have been proposed to address it.

The most straightforward one is based on an analysis in the time domain, and computes the period by measuring the distance between two successive zero-crossings [14]. It is very sensitive to noise.

The two most common methods are autocorrelation [9] in the time domain and spectral analysis with Fourier transform [23] in the frequency domain. They are efficient on specific data, namely sinusoidal and stationary signals, in which the period remains constant. The short-time Fourier [2] and wavelet [18] transforms are more sophisticated methods in the time-frequency domain able to deal with non stationary signals. However, the former needs a window size parameter to be efficient, and the latter is non parametric but very sensitive to time shifts, which is a bias to be avoided in the context of a high-level linguistic interpretation of the data.

Statistical methods have also been proposed, applying to the specific case where the data is sinusoidal with a Gaussian noise [7].

Lastly, cross-domain approaches have been developed, adding further complexity: in [3], as already mentioned, a fast Fourier transform is used on top of cyclic association rule extraction to build a list of candidate periods. In [26], both autocorrelation and a periodogram are used.

3 Evaluation of Periodic Protoforms

The principle of the Detection of Periodic Events methodology (DPE) [21] relies on the assumption that if a dataset contains regularly spaced high and low value groups of approximately constant size, then it is periodic. This assumption guides the truth evaluation of sentences of the form “Every p , values are high”. DPE is a modular methodology which can be seen as a general framework to evaluate periodic protoforms. This section describes its general architecture of the DPE methodology as well as its instantiation proposed in [21]. Section 4 presents a new watershed based method for its clustering step.

3.1 Input and Output

The input dataset, denoted X , is temporal and contains N normalised values (x_i), i.e. $X = \{x_i, i = 1, \dots, N\}$ such that $\forall i, x_i \in [0, 1]$. The data are considered to be regularly sampled, i.e. at date $t_i = t_1 + (i - 1) \times \Delta t$ where t_1 is the initial measurement time and Δt is the sampling rate.

The outputs of the DPE methodology are a periodicity degree π , a candidate period p_c and a natural language sentence. The periodicity degree π indicates the extent to which the dataset is periodic: 1 means it is absolutely periodic and the value decreases as the dataset is less periodic. The sentence is a linguistic description of the period found in the dataset, designed for human understanding. It has the form “ M every p unit, the data take high values”, where M is an adverb as “roughly”, “exactly”, “approximately”, p is the approximate value based on the candidate period p_c , and *unit* is a unit considered the most appropriate to express the period [21].

3.2 Architecture

The DPE methodology works in four steps : first, it clusters the data into groups of successive high or low values, second it computes the regularity of the group sizes and the periodicity degree, third it computes a candidate period and finally it returns a natural language sentence. In the following, more details are given regarding these 4 steps.

High and Low Value Detection The first step of DPE aims at detecting groups of high/low consecutive values. To this aim, a prediction function g returning the group type (H or L) of x_i is defined. Successive values classified H are gathered in high value groups, and conversely for low values.

A baseline function g_{BL} relies on a user-defined threshold t_{value} to distinguish high and low values :

$$g_{BL}(x_i) = \begin{cases} H & \text{if } x_i > t_{value} \\ L & \text{otherwise} \end{cases} \quad (1)$$

A function g_{ES} exploiting mathematical morphology tools is proposed in [21], based on the erosion score es defined as:

$$x_i^0 = x_i \quad x_i^j = \min(x_{i-1}^{j-1}, x_i^{j-1}, x_{i+1}^{j-1}) \quad es_i = \sum_{j=1}^z x_i^j$$

where z is the smallest integer such that $\forall i = 1 \dots N, x_i^z = 0$. This erosion score transform, classically used to identify the skeleton of a shape, has the following characteristics: high x_i in high regions have high es , low x_i in high regions have quite a high es , isolated high x_i in low regions have low es . Thus, erosion scores provide an automatic adaptation to the data level.

Computing the erosion score on the data complement $\bar{X}$ where $\bar{x}_i = 1 - x_i$ allow to symmetrically identify low regions. We propose the prediction function:

$$g_{ES}(x_i) = \begin{cases} H & \text{if } es_i > \bar{es}_i \\ L & \text{otherwise} \end{cases} \quad (2)$$

Groups are defined as successive values of the same type as returned by g .

Periodicity computing The second step of DPE consists in evaluating the regularity of the sizes of the high and low value groups. If these sizes are regular, then the dataset is considered periodic according to the assumption defined at the beginning of this section.

First, the size of each group is computed, setting $s_j^H = |G_j^H|$ for the j^{th} high value group and $s_j^L = |G_j^L|$ for the j^{th} low value group. Experiments with fuzzy cardinalities have showed no significant difference [21].

The regularity ρ is then determined for high and low value groups based on the average value μ and the deviation d of their size (see [21] for justification):

$$\rho = 1 - \min\left(\frac{d}{\mu}, 1\right) \quad \mu = \frac{1}{n} \sum_{j=1}^n s_j \quad d = \frac{1}{n} \sum_{j=1}^n |s_j - \mu| \quad (3)$$

both for high and low value groups n denotes the number of groups. The size dispersion is thus measured using the coefficient of variation $CV = d/\mu$: d is more robust to noise than standard deviation and the quotient with μ makes it relative and allows to adapt to the value level.

Finally, with the regularities of high value groups ρ^H and low value groups ρ^L , the periodicity degree π is returned as their average, i.e. $\pi = (\rho^H + \rho^L)/2$.

Candidate Period Computation For a perfectly regular phenomenon, the period is defined as the time elapsed between two occurrences of an event, in this paper, “high value”. Therefore the candidate period p_c is approximated as the sum of the average size of high and low value groups, i.e. $p_c = \mu^H + \mu^L$. p_c is relevant only if π is high enough, i.e. if the dataset is considered as periodic.

Linguistic Rendering The last step yields a linguistic periodic summary of the form “ M every p *unit*, the data take high values”, significant only if π is high. As described in [21], the *unit* used to describe the data is calculated first on a set of units entered as prior knowledge. Then the period p_c is rounded in order to make it more natural for a human being. Lastly, an adverb is chosen based on the approximation error between the computed and the rounded value.

4 Watershed Based Method

We propose a new method to identify high and low value groups based on another Mathematical Morphology tool, namely watershed. It can be seen as a variant of g_{BL} where the threshold is automatically derived from the data: it reduces the required expert knowledge and automatically adapts to the data.

4.1 Principle

Watershed in Mathematical Morphology has been introduced in [4] to perform 2D image segmentation. Its underlying intuition comes from geography: the image greyscale levels are seen as a topographic relief which is flooded by water. Watersheds are the divide lines of the domains of attraction of rain falling over the region. An efficient implementation has been proposed in [25] based on an immersion process analogy: as illustrated in Fig. 1, when the level of water rises, basins appear. When it rises more, new ones are created while others merge. At the end, all basins merge into a single one.

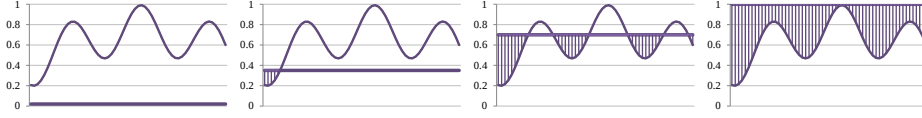


Figure 1. Illustration of the immersion of a dataset for the watershed calculation

We propose to apply watershed to detect groups, defining them from the identified basins for a given water line: low value groups are defined as the basins, i.e. consecutive values below the divide line, and high value groups as consecutive values above the divide line.

Furthermore, we propose to base the identification of the relevant water line, i.e. the threshold to separate high and low value groups, on the evolution of the basin structure. Indeed, the desired threshold should lead to a group identification that is robust to small local noise, i.e. making it a little greater or lower should not modify the number of identified groups. As water rises, basins appear (resp. disappear), when a gap (resp. a peak) crosses the water line: the threshold should be located at a level where no peak and gap, that represent local noise, are present. As formalised below, we propose to set the water line at the middle of the largest interval separating two consecutive basin structure changes.

4.2 Implementation

The changes of the basin structure are easily identified as they correspond to local peaks and gaps, i.e. values resp. greater or lower than their direct neighbours (previous and next values): when a peak or gap is identified, its level is recorded as a level where a basin structure change occurs. This principle is formalised below after the description of the pre-processing steps applied to the data.

Preprocessing Before finding the peaks and gaps, 2 preprocessing steps are applied: first, a moving average on a window whose size w is chosen by expert knowledge is calculated to smooth the curve and to avoid oversegmentation.

Second the consecutive equal values are removed. Indeed, basin structure changes could occur in configurations named “plateaux” which are different from gaps and peaks. A plateau is a collection of consecutive equal points surrounded with points of lesser (convex plateaux) or greater (concave plateaux) values. So as to ease the structure change detection and once the data are smoothed, consecutive equal points are removed from the dataset, so that convex plateaux become peaks, and concave plateaux become gaps. Thus, all basin structure changes can be detected with a simple peak/gap analysis.

Processing With the preprocessed data W , the determination of the levels where the basin structure changes is done with a single scan to detect local peaks and gaps. The levels at which the changes occur are stored in L :

$$L = \{w_i \in W / (w_i > w_{i+1} \wedge w_i > w_{i-1}) \vee (w_i < w_{i+1} \wedge w_i < w_{i-1})\}$$

Then the adaptable watershed-based threshold t_W is computed as:

$$t_W = \frac{1}{2} (L_m + L_{m+1}) \quad m = \arg \max_{i \in \{1 \dots |L|-1\}} L_{i+1} - L_i \quad (4)$$

where L_j are the elements of L sorted in ascending order. Finally, the clustering function is defined as:

$$g_W(x_i) = \begin{cases} H & \text{if } x_i > t_W \\ L & \text{otherwise} \end{cases} \quad (5)$$

5 Experimental Results

This section presents results obtained with artificial data, to compare the baseline, erosion score and watershed methods defined by (1), (2) and (5).

5.1 Data Generation

The datasets are generated as noisy series of periodic shapes, either rectangles or sines. They are created as a succession of high and low value groups, of size p^H and p^L respectively, on which two types of noise are applied: the group size noise ν_s randomly modifies the size values p^H and p^L , the value noise ν_y changes the values taken by the data within the groups. In the first step, the sizes of the high and low value groups are randomly drawn, adding some noise to the ideal values p^H and p^L : p^* generally denoting one of these two values, the size of each group is defined as:

$$s_j = \lceil 1 + \text{sgn}(0.5 - \epsilon_1) \times \nu_s \times \epsilon_2 \rceil p^*$$

where ϵ_1 and ϵ_2 are uniform random variables $\mathcal{U}(0,1)$. This distribution randomly increases or decreases the reference group size, through the $\text{sgn}(0.5 - \epsilon_1)$ coefficient, in a proportion defined as $\nu_s \times \epsilon_2$. The size of a group thus varies between $(1 - \nu_s)p^*$ and $(1 + \nu_s)p^*$. Group sizes are generated until their cumulative sum reaches the total desired number of points N .

After the group sizes have been determined for the rectangle shape, if the j^{th} group spans from index a to b , X^* is set as:

$$\forall k \in \{a, \dots, b\} \quad x_k^* = \begin{cases} 1 & \text{if group } j \text{ is high} \\ 0 & \text{otherwise} \end{cases}$$

For the sine shape, if the j^{th} group spans from index a to b , X^* is set as:

$$\forall k \in \{a, \dots, b\} \quad x_k^* = \frac{1}{2} + \frac{\lambda}{2} \sin\left(\pi \frac{k-a}{b-a}\right) \quad \lambda = \begin{cases} 1 & \text{if group } j \text{ is high} \\ -1 & \text{otherwise} \end{cases}$$

This calculation for the sine shape creates a discontinuous break around 0.5. It does not seem to introduce biases in the results though.

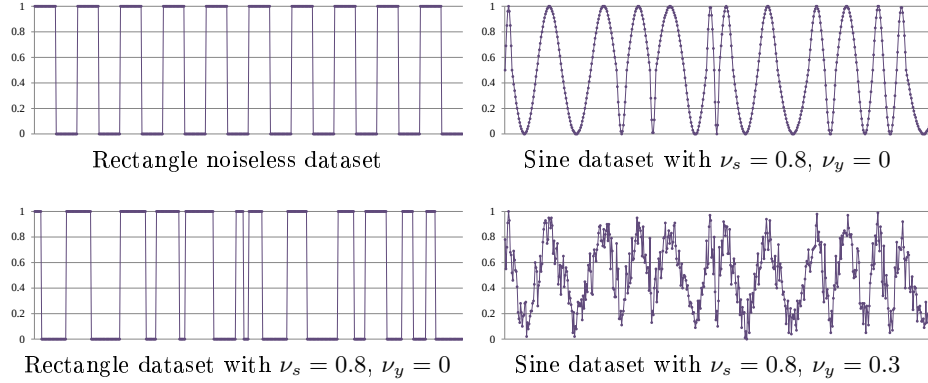


Figure 2. Four examples of generated datasets

In the third step the value noise ν_y is added to X^* leading to $\hat{X}$. The noise is applied downward for high value groups and upward for low value groups:

$$\hat{x}_i = \begin{cases} x_i^* - \nu_y \epsilon & \text{if } \hat{x}_i \text{ is in a high value group} \\ x_i^* + \nu_y \epsilon & \text{otherwise} \end{cases}$$

where ϵ is a uniform random variable $\mathcal{U}(0, 1)$.

Finally, the dataset X results from the normalisation of $\hat{X}$ to $[0, 1]$. Fig. 2 illustrates 4 examples of generated datasets.

5.2 Experimental protocol

As shown in Table 1, 16 test scenarios are implemented where the data series are generated with an increasing value noise ν_y and group size noise ν_s from 0 to 1 at a 0.05 pace (21 values) with a different combination of high / low value groups size, shape, group size noise and value noise.

The periodicity degree π , the candidate period p_c , the error in period evaluation Δp and the clustering accuracy Acc are computed with the 3 methods, baseline (BL), watershed (W) and erosion score (ES): Δp is defined as $\Delta p = |p_c - p| / p$. The period p to compare the candidate period with is computed as the sum of the average sizes of the two types of generated groups.

Table 1. The 16 test scenarios. The noise specified in the header is the constant one in the scenario, the not mentioned one is the increasing one.

Size (p^H/p^L)	Shape	$\nu_s = 0$	$\nu_y = 0$	$\nu_s = 0.5$	$\nu_y = 0.3$
Balanced (25/25)	Square	S1	S2	S9	S10
	Sine	S3	S4	S11	S12
Thin (10/40)	Square	S5	S6	S13	S14
	Sine	S7	S8	S15	S16

To compute Acc , the accuracy in the classification into high and low value groups, the labels are the group membership defined in the generation step. Acc is weighted so as to take into account the bias in the group size, for the “Thin (10/40)” scenarios.

5.3 Result interpretation

General results From the results obtained with the 16 scenarios and not detailed here, it appears that the noise type (group size or value) is the most important parameter: all methods exhibit similar behaviours for the 3 measures π , Δp and Acc for a given noise type. The shape is the second most important parameter, since differences appear between squares and sines, especially with increasing ν_y . Other parameters, as the size of the groups (balanced or thin) or the combination of noises, do not seem to bear an important influence.

The results also show that all 3 methods return a periodicity of 1 when the data has no noise and are decreasing functions of the noise parameters.

In the following, we focus on scenarios 5 to 8. Figure 3 illustrates the outcomes of the experiments. Nevertheless, the results mentioned below are valid for all considered scenarios.

Baseline approach The comparison between methods shows that the baseline curve is very sensitive to noise. Indeed, with squares (Fig. 3a, b), π falls sharply and Δp rises sharply as soon as ν_y reaches 0.5. This is due to the fact that from this level of noise, some points labelled as high have a value smaller than 0.5 and are classified as low. Since very few points are misclassified, accuracy is still high but small groups are created within larger ones, generating a high deviation in the group size, yielding poor periodicity degree and period evaluation precision.

This behaviour appears for lower values of ν_y with sines (Fig. 3g, h) since this kind of misclassification is possible as soon as $\nu_y > 0$.

Interestingly, the baseline Acc is always comparable to the one obtained with the other methods (Fig. 3c, f, i, l). Indeed, the phenomenon just described slightly affects the clustering accuracy. Moreover, as ν_y increases, the accuracy decreases in approximately the same amount for all methods, so BL remains comparable with the others. This is why the Acc measure is not very relevant here to choose a method, whereas Δp is much more discriminant.

Erosion score vs. watershed Generally speaking, erosion score is smoother than watershed in the sense that it varies less abruptly, ensuring steadiness in the evaluation. Moreover, it is generally more or equally precise in calculation of the period (11 scenarios out of 16) and in clustering accuracy (12 scenarios out of 16). Furthermore, the erosion score is also more precise over several experiences since it has a lower standard deviation than the watershed.

Regarding group size noise, watershed gives wrong period estimation when $\nu_g > 0.7$ (Fig. 3e). This is due to the fact that the moving average makes small peaks disappear. On the other hand, ES keeps these small peaks especially with

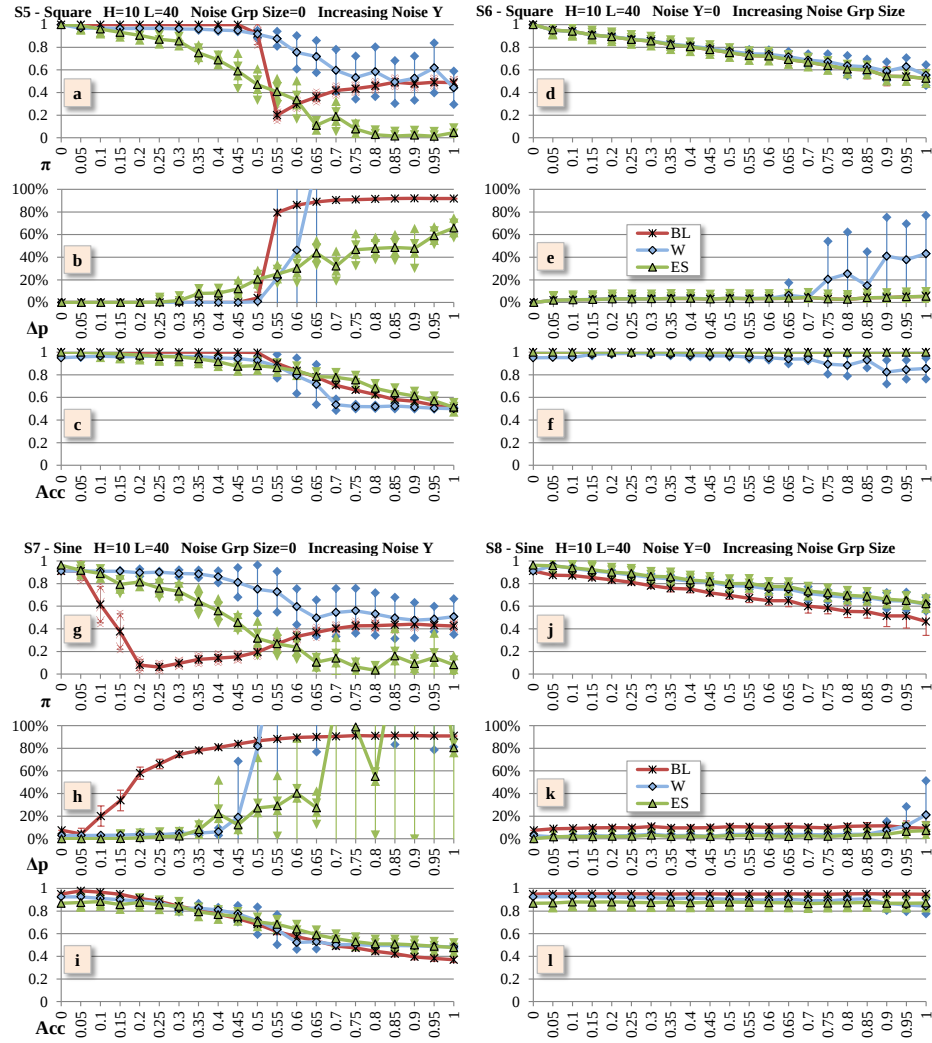


Figure 3. Mean and standard deviation for π , Δp and Acc for scenarios 5 to 8

3 graphs per scenario: top is periodicity degree, middle candidate period estimation error Δp , bottom clustering accuracy Acc .

3 curves per graph: the cross/red is the baseline BL, triangle/green is the erosion score ES and diamond/blue is the watershed W.

rectangles, which is a context where data are highly contrasted (0 or 1). With less contrasted data like sines, the difference is attenuated and all methods perform very well (Fig. 3k).

As for ν_y , the period computation becomes wrong as it increases, especially with sines (Fig. 3h). However, ES seems to be more robust to ν_y than watershed since the latter increases sharply from $\nu_y = 0.5$. This can be linked with the erosion score not using a constant threshold to cluster the data as opposed to the watershed method. Since the groups are processed individually with ES, a misclassification for one group does not necessarily propagate to the others, whereas a threshold not chosen appropriately with the watershed leads to misclassification throughout the dataset, resulting in a bad evaluation of the period.

6 Conclusion

A new watershed clustering method is proposed as an alternative to assess the relevance of linguistic expression of the form “ M every p unit, the data take high values” and tested within the framework of the DPE methodology [21]. Experimental results obtained with different shapes (rectangle and sine) and noises (value and group size) prove to be relevant. The DPE methodology is a good approach to classify the data in high and low value groups and to estimate the period of the dataset. The erosion score method seems more precise than the watershed one, due to its adaptable threshold for classification.

Future works aim at developing new fuzzy quantifiers as “from time to time”, “often”, “rarely”, and detect periodicity in sub-parts of the dataset, which both can be developed with the high and low values clustering in DPE. Another direction is the definition of a quality measure to compare the different methods, among themselves as well as to existing approaches.

References

1. Agrawal, R., Srikant, R.: Mining sequential patterns. In: Proc. of the 11th Int. Conf. on Data Engineering. pp. 3–14 (1995)
2. Allen, J., Rabiner, L.: A unified approach to short-time Fourier analysis and synthesis. Proc. of the IEEE 65(11), 1558–1564 (1977)
3. Berberidis, C., Vlahavas, I.P., Aref, W.G., Atallah, M.J., Elmagarmid, A.K.: On the Discovery of Weak Periodicities in Large Time Series. In: PKDD’02. pp. 51–61 (2002)
4. Beucher, S., Lantuejoul, C.: Use of watersheds in contour detection. In: Proc. of the Int. Workshop on image processing, real-time edge and motion detection/estimation (1979)
5. Bouchon-Meunier, B., Moyse, G.: Fuzzy Linguistic Summaries: Where Are We, Where Can We Go? In: CIFEr 2012. pp. 317–324 (2012)
6. Cariñena, P., Bugarín, A., Mucientes, M., Barro Ameneiro, S.: A language for expressing fuzzy temporal rules. Mathware & soft computing 7(2), 213–227 (2000)
7. Castillo, I., Lévy-Leduc, C., Matias, C.: Exact adaptive estimation of the shape of a periodic function with unknown period corrupted by white noise. Mathematical methods of statistics 15(2), 146–175 (2006)

8. Fu, T.C.: A review on time series data mining. *Engineering Applications of Artificial Intelligence* 24(1), 164–181 (2011)
9. Gerhard, D.: Pitch Extraction and Fundamental Frequency: History and Current Techniques. Tech. rep., University of Regina (2003)
10. Kacprzyk, J., Wilbik, A., Zadrozny, S.: Linguistic summarization of time series using a fuzzy quantifier driven aggregation. *Fuzzy Sets and Systems* 159(12), 1485–1499 (2008)
11. Kacprzyk, J., Yager, R.R.: "Softer" optimization and control models via fuzzy linguistic quantifiers. *Information Sciences* 34(2), 157–178 (1984)
12. Kacprzyk, J., Zadrozny, S.: Protoforms of Linguistic Data Summaries: Towards More General Natural-Language-Based Data Mining Tools. In: Abraham, A., Ruizdel Solar, J., Koeppen, M. (eds.) *Soft computing systems*. pp. 417–425 (2002)
13. Kacprzyk, J., Zadrozny, S.: Computing With Words Is an Implementable Paradigm: Fuzzy Queries, Linguistic Data Summaries, and Natural-Language Generation. *IEEE Transactions on Fuzzy Systems* 18(3), 461–472 (2010)
14. Kedem, B.: Spectral analysis and discrimination by zero-crossings. *Proc. of the IEEE* 74(11), 1477–1493 (1986)
15. Laxman, S., Sastry, P.S.: A survey of temporal data mining. *Sadhana* 31(2), 173–198 (2006)
16. Lee, C.H., Chen, M.S., Lin, C.R.: Progressive partition miner: An efficient algorithm for mining general temporal association rules. *IEEE Transactions on Knowledge and Data Engineering* 15(4), 1004–1017 (2003)
17. Ma, S., Hellerstein, J.L.: Mining partially periodic event patterns with unknown periods. In: *Proc. of the 17th Int. Conf. on Data Engineering*. pp. 205–214. IEEE Comput. Soc (2001)
18. Mallat, S.: A theory for multiresolution signal decomposition: the wavelet representation. *IEEE Trans. on PAMI* 11(7), 674–693 (1989)
19. Mannila, H., Toivonen, H., Inkeri Verkamo, A.: Discovery of Frequent Episodes in Event Sequences. *Data Mining and Knowledge Discovery* 1(3), 259–289 (1997)
20. Méger, N., Rigotti, C.: Constraint-based mining of episode rules and optimal window sizes. In: *Proc. of the 8th European Conf. on Principles and Practice of Knowledge Discovery in Databases*. pp. 313–324. Springer-Verlag (2004)
21. Moyse, G., Lesot, M.J., Bouchon-Meunier, B.: Linguistic summaries for periodicity detection based on mathematical morphology. In: *IEEE SSCT'13*. pp. 106–113 (2013)
22. Ozden, B., Ramaswamy, S., Silberschatz, A.: Cyclic association rules. In: *Proc. of the 14th Int. Conf. on Data Engineering*. pp. 412–421. IEEE Comput. Soc (1998)
23. Palmer, L.: Coarse frequency estimation using the discrete Fourier transform. *IEEE Transactions on Information Theory* 20(1), 104–109 (1974)
24. Serra, J.: Introduction to mathematical morphology. *Computer Vision, Graphics, and Image Processing* 35(3), 283–305 (1986)
25. Vincent, L., Soille, P.: Watersheds in digital spaces: an efficient algorithm based on immersion simulations. *IEEE Trans. on PAMI* 13(6), 583–598 (1991)
26. Vlachos, M., Yu, P.S., Castelli, V.: On periodicity detection and structural periodic similarity. *Proc. of the 5th SIAM Int. Conf. on Data Mining* 119, 449 (2005)
27. Yager, R.R.: A new approach to the summarization of data. *Information Sciences* 28(1), 69–86 (1982)
28. Zadeh, L.A.: A computational approach to fuzzy quantifiers in natural languages. *Computers & Mathematics with Applications* 9(1), 149–184 (1983)