



Classifying pollutant flush signals in stormwater using functional data analysis on TSS MV curves

Ditte Marie Reinholdt Jensen^{a,b,c,*}, Santiago Sandoval^{d,e}, Jean-Baptiste Aubin^d, Jean-Luc Bertrand-Krajewski^d, Li Xuyong^b, Peter Steen Mikkelsen^a, Luca Vezzaro^a

^a Department of Environmental and Resource Engineering, Technical University of Denmark (DTU), Bygningstorvet, Bygning 115, 2800 Kongens Lyngby, Denmark

^b State Key Laboratory of Urban and Regional Ecology, Research Center for Eco-Environmental Sciences (RCEES), Chinese Academy of Sciences (CAS), 18 Shuangqing Road, Beijing 100085, China

^c Sino-Danish Center for Education and Research (SDC), Aarhus, Denmark and University of Chinese Academy of Sciences (UCAS), China

^d University of Lyon, INSA Lyon, DEEP, EA 7429, F-69621 Villeurbanne cedex, France

^e University of Applied Sciences and Arts of Western Switzerland (HES-SO), HEIA-FR, ITEC, Boulevard de Pérolles 80, 1700 Fribourg, Switzerland

ARTICLE INFO

Keywords:

Pollutant flush
Separate sewer systems
MV curves
Clustering
Urban stormwater management

ABSTRACT

Pollution levels in stormwater vary significantly during rain events, with pollutant flushes carrying a major fraction of an event pollutant load in a short period. Understanding these flushes is thus essential for stormwater management. However, current studies mainly focus on describing the first flush or are limited by predetermined flush categories. This study provides a new perspective on the topic by applying data-driven approaches to categorise Mass Volume (MV) curves for TSS into distinct classes of flush tailored to specific monitoring location. Functional Data Analysis (FDA) was used to investigate the dynamics of MV curves in two large data sets, consisting of 343 measured events and 915 modelled events, respectively. Potential links between classes of MV curves and combinations of rain characteristics were explored through a priori clustering. This yielded correct class assignments for 23–63% of the events using different combinations of MV curve clustering and rainfall characteristics. This suggests that while global rainfall characteristics influence flush, they are not sufficient as sole explanatory variables of different flush phenomena, and additional explanatory variables are needed to assign MV curves into classes with a predictive power that is suitable for e.g. design of stormwater control measures. Our results highlight the great potential of the FDA methodology as a new approach for classifying, describing, and understanding pollutant flush signals in stormwater.

1. Introduction

Pollutants carried by urban runoff in separate sewer systems pose an environmental threat to the natural receiving water bodies (Brudler et al., 2019; Eriksson et al., 2011; Göbel et al., 2007; Koziel et al., 2019; Walsh et al., 2005; Zgheib et al., 2011) as well as to human health, if e.g. runoff is used as source of potable water (Kus et al., 2010; Ma et al., 2019). To enable management actions aiming to reduce and mitigate this threat it is thus critical to gain knowledge of expected pollutant levels and of their spatial and temporal variations.

Pollutants in stormwater runoff stem from multiple release processes in the catchment: dry and wet atmospheric deposition, surface weathering, leaching from building materials, and release from anthropogenic

sources, such as the use of specific chemicals (pesticides, spills, etc.) or vehicular activity (Deletic and Orr, 2005; Eriksson et al., 2005; Göbel et al., 2007; Müller et al., 2020). During rain events, these pollutants are washed off and transported away by rainfall and by the generated runoff. The resulting runoff quality thus depends on the factors driving the pollutant accumulation, the mobilisation and the transport mechanisms, resulting in large inter and intra-events variations in both concentrations and loads (Bertrand-Krajewski et al., 1998; Lee et al., 2002; Qin et al., 2016).

Historically, there has been a vast interest in describing pollutant flushes during a rain event, with great emphasis often being put on the *first flush*, i.e. pollutant flushes carrying out the majority of the event load during the first part of the event. The occurrence of first flush have

* Corresponding author.

E-mail addresses: dije@env.dtu.dk (D.M. Reinholdt Jensen), santiago.sandoval@hefr.ch (S. Sandoval), jean-baptiste.aubin@insa-lyon.fr (J.-B. Aubin), xyli@rcees.ac.cn (L. Xuyong), psmi@env.dtu.dk (P.S. Mikkelsen), luve@env.dtu.dk (L. Vezzaro).

<https://doi.org/10.1016/j.watres.2022.118394>

Received 24 November 2021; Received in revised form 16 March 2022; Accepted 2 April 2022

Available online 4 April 2022

0043-1354/© 2022 The Authors. Published by Elsevier Ltd. This is an open access article under the CC BY license (<http://creativecommons.org/licenses/by/4.0/>).

been studied by a great number of studies, proposing different definition schemes, and developing different classification to describe this phenomenon (Jensen, 2022).

The increasing amount of available data on stormwater runoff quality has enabled the application of data-driven and statistical methods for exploring these datasets and gaining knowledge on pollutant flushes. For example, Deletic (1998) applied linear multi-regression analysis to search for correlations between the flush in the first 20% of the runoff and combinations of rainfall and runoff characteristics. Both Bach et al. (2010) and Sun et al. (2015) used the Wilcoxon Rank Sum test to group pieces of runoff events with similar concentration levels to determine the catchment-specific volume carrying the first flush. Li et al. (2012) applied redundancy analysis to quantify the influence of rain and catchment characteristics on the variation in the strength of the first flush. Li et al. (2015) used Pearson correlation analysis to connect Event Mean Concentrations (EMCs) and the strength of flush in the first 40% of the runoff to catchment and rain characteristics. Finally, Perera et al. (2019) applied a random forest, tree-based method in order to identify nonlinear relationships between rain and catchment characteristics and first flush.

However, a common denominator in studies of flush dynamics (including the data-driven studies mentioned above) is their sole focus on describing the first flush or their limitation to predetermined flush categories (Jensen, 2022). This exclusive focus might neglect other site-specific (and perhaps more relevant) tendencies and patterns, disregarding the span of variation in inter- and intra-event pollutant distributions that deviate from the general behaviour defined in literature. Indeed, there is a gap in the use of more exploratory approaches, with a broader scope than (and not limited to) the first flush. Such approaches can support stormwater planners in developing robust practices for managing pollutants in urban runoff by considering all possible flush patterns in a site and their impacts on the planned stormwater control measure.

Functional data analysis (FDA) is a strong tool for revealing dynamics of underlying processes in repeated observations (Wang et al., 2016), and it has the added benefit of not relying on approximations to a specific curve shape. As such, FDA has been employed for different purposes within the field of water management, e.g. for analysing spatial and temporal variabilities of rainfall data in order to identify basis functions of rainfall pattern at different stations or seasons (Suhaila and Yusop, 2017), for detecting anomalies in flow within water supply networks (Millán-Roures et al., 2018) or classifying streamflow hydrographs (Ternynck et al., 2016). These examples show how FDA was applied for classifying functional data, determining main trends or detecting outliers. To the best of our knowledge, this methodology has never been applied for analysing stormwater pollutant flush dynamics - possibly because it requires the availability of rather large data sets.

This study aims to provide a better understanding of pollutant flushes in separate sewer systems by using unsupervised analysis tools based on FDA to (i) explore and categorise distinct Mass Volume (MV) curve shapes and (ii) describe their frequency of occurrence. Furthermore, we (iii) assess the potential for linking the resulting flush classes to combinations of rain variables. The study is undertaken using large data sets of events from both monitored and modelled case studies.

2. Materials and methods

2.1. Theory of MV curves and flush drivers

A widely applied approach to describe pollutant flushes during an event is based on MV curves, which have been extensively investigated in scientific literature (e.g. Bertrand-Krajewski et al., 1998; Chow and Yusop, 2014; Deletic and Maksimovic, 1998; Kong et al., 2021; Métadier and Bertrand-Krajewski, 2012; Perera et al., 2021). MV curves present the relative cumulative mass (M) as a function of the relative cumulative volume (V), enabling the analysis of flush signals (see Fig. 1). The first

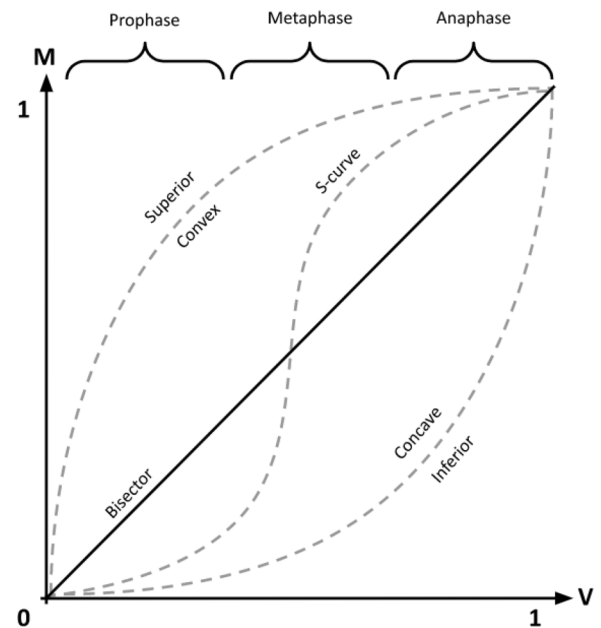


Fig. 1. Theoretical representation of MV curves and introduction of terminology. Inspired by Qin et al. (2016) and Métadier and Bertrand-Krajewski (2012).

flush takes place if the majority of the mass is carried in the prophase of the runoff volume, and its strength is described by the convexity of the (superior) MV curve. No flush occurs if the MV curve follows the bisector, while a second flush (or middle flush) occurs when the MV curve follows an s-shape, with the majority of the mass discharged in the metaphase. If the MV curve is concave (inferior), it indicates that the majority of the mass is discharged in the anaphase, resulting in third flush (or late flush). The MV curve approach thus enables the comparison of flush signals across events with different mass, volume, and duration. However, several flushes may occur during a single event, resulting in a large variety of possible MV curve shapes. MV curves are often fit to polynomial approximations in order to conduct data analysis (e.g. Métadier and Bertrand-Krajewski, 2012). This may be a good approximation for superior and inferior curves, but it is a poor estimation for s-curves or reverse s-curves.

While the event normalisation allows a comparison of dynamics across events that may otherwise be hard to compare, other relevant event characteristics (magnitude, duration, return period, etc.) are lost. Therefore, if relevant for e.g. design purpose, events should be pre-selected based on these characteristics before the MV curve normalisation and analysis.

Although several studies investigating the drivers of flush signals are based on MV curves, it is often hard to compare their results due to the use of different flush definitions, focus on different pollutants, and different local conditions. Indeed, Deletic (1998) noticed that the influential rainfall and runoff characteristics might differ for catchments of very similar characteristics. Perera et al. (2019) found that rain-related variables were ranked highest in terms of influence on first flush and suggested that their combination might yield better results. While the occurrence of first flush has been found to be affected by the type of pollutant (particulate or dissolved), the system (combined or separate), the catchment characteristics (land use, slope, size, etc.), the sampling technique, and the flush definition (Jensen, 2022), rain characteristics are among the most typically explored drivers - in part as they are more frequently accessible. One explanation could be that rain characteristics are assumed to be the major driver of inter-event variability, while other drivers are often assumed to be constant or to have little influence at the temporal scale.

A great disagreement is found in literature on the influence of the Antecedent Dry Weather Period (ADWP), also called Antecedent Dry

Weather Days (ADWD) or Antecedent Dry Days (ADD) (e.g. Chow and Yusop, 2014; Deletic and Maksimovic, 1998; Li et al., 2012; Schriewer et al., 2008). This is notable, as this variable constitutes the backbone of build-up modules in several accumulation-washoff models.

Several studies stressed the influence of rain intensity (e.g. Chow and Yusop, 2014; Deletic and Maksimovic, 1998; Li et al., 2012). However, Athanasiadis et al. (2010) found that intensity did not affect first flush effects of Cu in roof runoff, and Schriewer et al. (2008) found that Zn in roof runoff is correlated with lower rain intensities. Furthermore, Kong et al. (2021) found that first flush effect depended on the timing on the maximum intensity, and similarly, mixed interpretations can be found for the influence of total rain depth of an event and its duration (Athanasiadis et al., 2010; Chow and Yusop, 2014; Kang et al., 2008; Li et al., 2015; Perera et al., 2019).

As studies on influential flush factors often point in different directions, it is likely that the governing explanatory variables are case-specific. This stresses the need for approaches capable of evaluating flush signals systematically, rather than general definitions that may not apply under local conditions.

2.2. Case studies

In order to capture several different flush signals it is necessary to utilise an extensive data set from a specific location. Therefore, the analysis was carried out on a large data set of monitored Total Suspended Solids (TSS) data from Chassieu (France). As large data sets of TSS data are seldom available (with Chassieu being one of the few examples worldwide), these were supplemented by a set of modelled data from Albertslund (Denmark), obtained using a state-of-the-art accumulation-washoff model. While inherently uncertain due to the model simplification of reality, the modelled data have the additional benefit of being less affected by measurements noise and they are thereby assumed easier to analyse. For both data sets TSS concentration time series are used as indicator for pollutant levels.

2.2.1. Chassieu data set (monitored)

The Chassieu catchment covers an industrial area in the outskirts of Lyon (France) drained by a separate sewer system. Continuous monitoring of rain (6-minute time steps), flow and turbidity (2-minute time steps) was carried out from 2004 to 2011. The data set is described in Métadier and Bertrand-Krajewski (2012) and Sun et al. (2015), while the conversion from turbidity to TSS is reported in Métadier and Bertrand-Krajewski (2011). The data set includes 716 events, but not all of these contain complete measurement series (see Table 1). For further details, we refer to the previous works conducted on this data set (e.g. Métadier and Bertrand-Krajewski, 2011; 2012; Sandoval et al., 2018; Sun et al., 2015).

2.2.2. Albertslund data set (modelled)

The Albertslund catchment covers a residential/industrial area in the outskirts of Copenhagen (Denmark) drained by a separate sewer system. Continuous monitoring of rain, flow and turbidity was carried out in 2010–2011. These measurements were used to calibrate a dynamic

Table 1

Main characteristics of the two analysed case studies. For more details please refer to Métadier and Bertrand-Krajewski (2012) and Vezzaro et al. (2015).

	A_{imp}	Events*	Rain depth _{avr}	Rain du _r _{avr}	Rain int _{avr}	TSS SMC**
Chassieu	133 ha	343	8.49 mm	6.45 h	1.77 mm/h	119 g/m ³
Albertslund	94.7 ha	915	5.33 mm	5.37 h	1.26 mm/h	63 g/m ³

*The presented values are after step 1a in the methodology has been completed.

**SMC = Site Mean Concentration.

stormwater pollution model, based on the state-of-the-art accumulation-washoff process (Vezzaro et al., 2015). The model parameters were estimated by using the pseudo-bayesian, uncertainty-based methodology described in Vezzaro et al. (2013). To propagate model uncertainties, a total of 1,000 parameter sets were sampled from the behavioural parameters identified in (Vezzaro et al., 2015). These were used to simulate a 10-year period (using rainfall data covering the period 1994–2004), resulting in a total of 1,094 discharge events. The characteristics of each event (volume, load, etc.) were calculated on the median of the 1,000 simulations (see Table 1).

2.3. Data analysis

This study utilises a data-driven and exploratory approach based on techniques from Functional Data Analysis (FDA), which allow for the identification of multiple categories and types of flush signals without relying on predetermined definition schemes. An overview of the performed steps for both case studies is given in Fig. 2 and further elaborated in the following subsections:

1. Data preparation: a validation of the events was carried out, selecting those to be analysed, preparing MV curves, and extracting rain characteristics for the further steps.
2. MV curve analysis: this step explored the dynamics of MV curves and divided them into classes based on their shape.
3. Rain variables analysis: non-essential variables were eliminated before the combined analysis.
4. Combined analysis: this step explored the direct relationship between MV classes and rain variables and indirect relationships between MV classes and combinations of rain variables.

2.3.1. Data preparation

Events in the monitored data set (Chassieu) that contained missing values of runoff flow or TSS concentration for more than 10 continuous minutes were discarded. Missing values for less than 10 continuous minutes were linearly interpolated. The modelled data set (Albertslund) included small rainfall events that did not generate runoff. Therefore, events with a rain duration of shorter than 10 minutes were eliminated, as generally these events had very low rain depth.

For each event in the data set, MV curves were generated using TSS and flow data. To ensure that all MV curves had equal discretisation (a prerequisite for the used FDA algorithms), they were rescaled using interpolation to ensure a discretisation of 101 steps per curve (corresponding to a value of 0%, 1%, ..., 100% of the duration of the runoff event).

An overview of the investigated rain characteristics is provided in Table 2, which includes both variables found to be significant by previous studies and additional new variables. For the Chassieu data set,

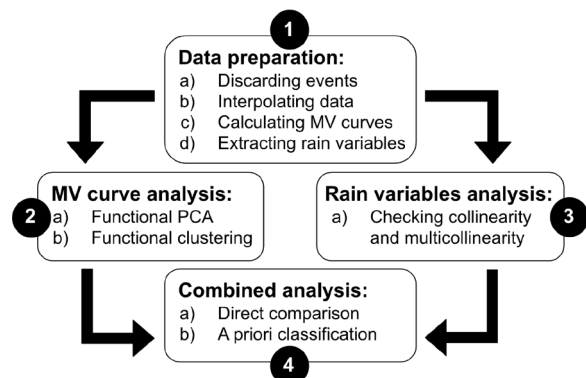


Fig. 2. Overview of applied methodologies and flow of approach.

Table 2

Overview of the investigated rain characteristic variables and their abbreviations. Variables marked in bold are those that passed through step 3 of the methodology (based on the qualitative assessment of plots found in Supplementary Material).

Acronym	Description	Unit
DUR	Total duration of rainfall	min
AVR_INT	Average rainfall intensity	$\mu\text{m/s}$
MAX_INT	The maximum 6-minute intensity in the rainfall event	$\mu\text{m/s}$
MAX_STEP	The number of the time step where the MAX_INT begins	number
MAX_STEP_REL	MAX_STEP divided by total number of time steps in the rainfall event	fraction
MAX_INT_50FIRST	Maximum intensity (6-minute steps) found in the first 50% of the total rain depth	$\mu\text{m/s}$
MAX_INT_50LAST	Maximum intensity (6-minute steps) found in the last 50% of the total rain depth	$\mu\text{m/s}$
MAX_INT_50REL	MAX_INT_50FIRST divided by MAX_INT_50LAST	fraction
DEPTH	Total depth of the rainfall	mm
NO_RAIN	The amount of minutes in the event without any rain divided by the total duration	fraction
ADWP	Antecedent dry weather period	hours
AR2	Antecedent rainfall 2 days from the beginning of the event	mm
AR5	Antecedent rainfall 5 days from the beginning of the event	mm
AR7	Antecedent rainfall 7 days from the beginning of the event	mm
AR14	Antecedent rainfall 14 days from the beginning of the event	mm
Max_int_last	Maximum 6-minute intensity in the last, preceding rain event	$\mu\text{m/s}$
Month	The month the event started in (for seasonality)	no.

rain data was only available for the events in the database (i.e. there were no data for the periods between the recorded events), so the ADWP was estimated based on the 716 available events. Also, for both the Chassieu and the Albertslund data sets, there was no information about rain preceding the first event in the database, so ADWP and values of antecedent rain in the last 2-14 days (AR2-14) were set to zero.

2.3.2. MV curve analysis

A functional principal components analysis (fPCA) was run to explore the shape and dynamics of the MV curves, followed by functional clustering to divide the MV curves into classes. fPCA is the most common first step of any FDA, as it reduces the dimensionality of the longitudinal data and creates a foundation for many other FDA methods (Chen et al., 2017), in addition to offering insight in itself. It is based on an eigenvalue analysis of the underlying covariance operators of the observed process. fPCA was used to generate eigenfunctions for the functional data until a high percentage of the variance (99.9%) was described. Each eigenfunction has a corresponding list of eigenvalues with length matching the number of events in the data set. This resembles the Principal Component (PC) scores for that particular PC index. Both the fPCA and the functional clustering were carried out using the 'fdapace' package (Carroll et al., 2020), which is an open-source R-package for FDA and empirical dynamics (an equivalent 'PACE' package is also available for Matlab: <http://www.stat.ucdavis.edu/PACE/>).

The functional clustering was carried out to identify natural MV profile classes in an unsupervised manner through strict partitioning clustering, meaning that each object belongs to exactly one cluster. The fdapace package offers two options for functional clustering: the Expectation Maximization Cluster (EMCluster) algorithm and the k-Centers Functional Clustering (kCFC) algorithm. EMCluster uses the fPC scores previously calculated through fdapace (Carroll et al., 2020) and creates multivariate Gaussian distributions for each cluster (Chen and Maitra, 2015). kCFC looks at both the mean and modes of the variation differentials between clusters by predicting the cluster membership

through an iterative covariance updating scheme. Here, local linear smoothing is applied, covariance for each curve is calculated separately and aggregated to estimate the smoothed covariance, followed by eigenanalysis and calculation of fPC scores using conditional expectation (Chiou and Li, 2007). Both algorithms were tested to find the best way to divide the data set into a natural number of clusters. A comparison was made through visual inspection of the resulting clusters and by looking at the Sum of Squared Errors (SSE), i.e. the summed and squared mean distance from each fPC score (across all dimensions needed to reach 99.9% explained variance) to the centre of its associated cluster.

2.3.3. Rain variables analysis

In order to prepare the rain variables for the PCA and cluster analysis, they were first checked for collinearity, i.e. whether any of the variables were linear predictors of each other. While PCA can in itself remove collinearity, PC scores will be affected by the redundant features included in the analysis through collinear variables. Clustering collinear variables can indeed place additional weight on a feature because it is represented in multiple variables. Although removing collinear variables is usually considered not to affect k-means clustering, it has the extra benefit of speeding up calculations and achieving a more parsimonious model.

The collinearity was evaluated by looking at correlation plots of all the rain variables. Multivariate collinearity was evaluated by running a linear regression (with the fPC scores as response variables and rain variables as explanatory variables), and evaluating the Variance Inflation Factor (VIF) of the resulting regression.

2.3.4. Combined analysis

Firstly, a direct comparison of the MV curves and rain variables was conducted by looking at box plots of rain variable values for each cluster from the functional clustering analysis. Furthermore, a Kruskal-Wallis test was carried out to ascertain if significant difference had been obtained between the clusters for each rain variable (Kruskal and Wallis, 1952). It is a nonparametric, global test across all classes at once, consistent with the data-driven approach to cluster-analysis followed in this paper. The test uses a null hypothesis that the median (for each rain variable) is comparable across groupings (MV classes), which is rejected if any of the classes stand out with a significantly different value. If the test shows that some of the MV classes provide values of rain characteristics significantly different from the rest, this indicates a potential for using rain variables as explanatory factors of the MV classes.

Secondly, a clustering scheme was set up to combine the rain variables, investigating if their joint variation can be linked to the MV categories. The analysis first looked at using a single rain variable, and then an increasing number of variables were combined. Before clustering, the variables were standardised and ran through a PCA to reduce dimensionality and increase the efficiency of the analysis. Clustering was carried out using a standard k-means approach. The MV categories were used as an a priori classification for the clustering of the rain variables, and a contingency table (as shown in Table 3) was produced for each

Table 3

Contingency table for comparing partition X of R clusters with partition Y of C clusters. Modified from Hubert and Arabie (1985).

	Y_1	Y_2	$\dots$	Y_C	Sums
X_1	n_{11}	n_{12}	$\dots$	n_{1C}	a_1
X_2	n_{21}	n_{22}	$\dots$	n_{2C}	a_2
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
X_R	n_{R1}	n_{R2}	$\vdots$	n_{RC}	a_R
Sums	b_1	b_2	$\vdots$	b_C	n

combination of rain variables. To evaluate the performance, we used the two indicators applied by [Chiou and Li \(2007\)](#) for comparison of two partitionings (cRate and aRand):

- **cRate**: the maximal possible ratio of correctly classified objects (bounded by 0 and 1), calculated as the sum of the diagonal in the contingency matrix (see [Table 3](#)) divided by the total number of events, as the diagonal represents events placed in the same cluster by both partitions.
- **aRand**: a corrected version of the Rand index ([Rand, 1971](#)) proposed by [Hubert and Arabie \(1985\)](#), which describes the agreement between two partitions of objects (events in this study) by looking at the number of paired objects that are in the same or different groups in both partitions. It is calculated as in [Eq. 1](#) (with nomenclature referring to the contingency table in [Table 3](#)). Here, a comparison is made between clusters of partitions X and Y, with a_i as the total amount of events placed in the i^{th} cluster in X, b_j as the total amount of events placed in the j^{th} cluster in Y, n_{ij} as the amount of events from cluster i in the X partition placed in cluster j in the Y partition, and n as the total number of events. aRand has an upper bound of 1 and an expected value of 0, although negative values may occur if the degree of similarity between the two partitions is less than what would be expected in the case of two random partitionings.

$$\text{aRand} = \frac{\text{Index} - \text{Expected index}}{\text{Maximum index} - \text{Expected index}} = \frac{\sum_{i,j} \binom{n_{ij}}{2} - \left[\sum_i \binom{a_i}{2} \sum_j \binom{b_j}{2} \right] / \binom{n}{2}}{\frac{1}{2} \left[\sum_i \binom{a_i}{2} + \sum_j \binom{b_j}{2} \right] - \left[\sum_i \binom{a_i}{2} \sum_j \binom{b_j}{2} \right] / \binom{n}{2}} \quad (1)$$

3. Results and discussion

3.1. Data preparation

The event discarding step led to a reduction in the number of events

in each data set: from 716 to 343 events for the monitored data set (Chassieu) and from 1094 to 915 events for the modelled data (Albertslund). A summary of the validated events is given in [Table 1](#).

The MV curves for the Chassieu and Albertslund sites are shown in [Fig. 3](#). Both plots show a wide span of MV curves on both sides of the bisector rather than any single, general trend in curve profile. Interestingly, the accumulation-washoff model does not produce evident first flush signals (see [Fig. 3b](#)), suggesting that differences in input rainfall dynamics between events are a major driver behind inter- and intra-event variation in MV curves.

As expected, the MV curves show a greater variability in the measured Chassieu data set compared to the modelled Albertslund data. The simplification and structural limitations of the model structure (accumulation-washoff) clearly lumped the variety of processes taking place in the natural system. Although accumulation-washoff processes are often used to support the theory behind first flush hypothesis, the simulation results display that the dynamic behaviour of rainfall plays a major role in the pollutant mobilisation in the catchment and transport in the drainage system, and thereby in the presence of flush signals. The mean of the MV curves shows a slight tendency of a middle flush in Chassieu, while the mean in Albertslund is closer to the bisector.

The extraction of rain characteristics from the data set resulted in seventeen different rain variables (see [Table 2](#)).

3.2. MV curve analysis

More than 96% of the Chassieu variability and more than 99% of the Albertslund variability is summarised in the first 3 principal components of the fPCA (see [Fig. 4a](#) and [4 b](#)). The large amount of information included in such few components shows a high level of structure in the data set (i.e. the curves share many similarities). The first three

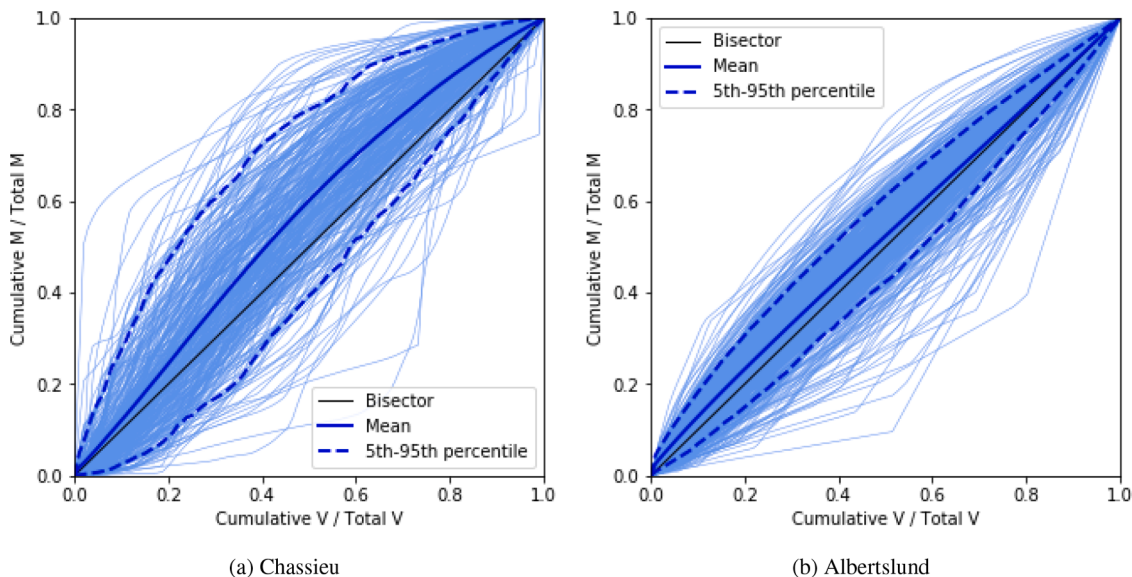


Fig. 3. Dimensionless MV curves for the two studied catchments: Chassieu (monitored) and Albertslund (modelled).