

Centre Eau Terre et Environnement

Développement d'approches statistiques multivariées pour les modèles de copules, applications en hydrologie

Par

Imene Ben Nasr

Thèse présentée pour l'obtention du grade de
Philosophiae doctor (Ph.D.) en sciences de l'eau

Jury d'évaluation

Présidente du jury et examinatrice interne	Sophie Duchesne INRS-ETE
Examineur externe	Jean-François Quessy Université de Québec à Trois-Rivières
Examineur externe	Taoufik Bouezmarni Université de Sherbrooke
Directeur de recherche	Fateh Chebana INRS-ETE

Résumé

Les études portant sur les évènements hydrologiques extrêmes sont d'une grande importance vu leurs nombreux impacts socio-économiques. Ces études reposent dans une large mesure sur la capacité à estimer adéquatement les risques associés à ces évènements. Dans ce cadre, l'analyse fréquentielle (*AF*) est une des approches les plus utilisées pour la modélisation et l'estimation des risques associés aux évènements extrêmes. Généralement, ces évènements sont caractérisés par plusieurs variables dépendantes, comme le volume, la pointe et la durée pour les crues. Par conséquent, l'analyse de chacune de ces variables séparément ne peut pas fournir une évaluation complète des risques et peut engendrer des pertes de vies humaines ou de biens associés à une sous-estimation, ou une augmentation des coûts des ouvrages hydrauliques associés à une surestimation. L'*AF* multivariée (*AFM*) permet de pallier ce problème en considérant simultanément ces variables.

L'*AF* classique est basée sur trois hypothèses à savoir l'homogénéité, la stationnarité et l'indépendance. Par ailleurs, différentes conditions non standards telles que la complexité topographique, les perturbations par les aménagements urbains et les changements climatiques peuvent influencer la réponse hydrologique. De telles conditions rendent la prédétermination des crues, par les méthodes classiques d'*AFM*, un exercice non efficace et mal adapté à de tels contextes. L'objectif de cette thèse consiste à proposer de nouvelles méthodes plus prometteuses pour l'analyse et la modélisation des variables hydrologiques à la fois dans un cadre multivarié et en l'absence de l'hypothèse d'homogénéité. Ces méthodes visent à contourner les limites de celles utilisées dans la littérature. Ces nouvelles méthodes, basées sur les copules, permettent une meilleure estimation des risques des extrêmes hydrologiques en tenant compte des interactions et

des dépendances entre les différentes variables. Par conséquent, les gestionnaires des ressources hydriques peuvent prendre des décisions éclairées et mieux adaptées au contexte actuel.

Précisément, nous nous intéressons dans une première partie à tester l'homogénéité des séries hydrologiques multivariées. Pour ce faire, nous avons proposé un nouveau test, basé sur les L-moments multivariés, capable de détecter la rupture dans la structure de dépendance des variables hydrologiques. Les résultats obtenus montrent la bonne performance du test proposé sous différents scénarios. Par la suite, dans le but de modéliser l'hétérogénéité des séries hydrologiques multivariées, un modèle de mélange de copules a été mis au point. Dans ce cadre, nous avons proposé une nouvelle méthode pour l'estimation des différents paramètres du modèle. La méthode proposée est basée sur l'utilisation des algorithmes génétiques. Sa performance est illustrée sur des séries simulées. Les résultats obtenus montrent que la méthode proposée est un outil efficace capable de fournir de bonnes estimations, en particulier dans un contexte hydrologique. Finalement, deux nouveaux tests d'adéquation pour les copules multiparamètres ont été développés. Le but de ces tests est de vérifier la qualité de l'ajustement de la copule aux données et améliorer ainsi l'estimation des risques associés aux événements extrêmes. Les résultats révèlent que les tests proposés performant bien dans le contexte hydrologique.

Mots-clés : Analyse fréquentielle; Homogénéité; Dépendance; Copule multiparamètre; Modèle Mixte; Estimation; Tests d'ajustement; L-moments multivariés; Crue; Maximum pseudo-vraisemblance; Méta-heuristique; Algorithmes génétiques; Maximisation-Espérance.

Abstract

Studies of hydrological extreme events are of great importance given their many socio-economic impacts. These studies rely on the ability to adequately estimate the associated risk. In this context, frequency analysis (*FA*) is one of the most widely used approaches for modeling and estimating the risks associated with extreme events. Generally, these events are characterized by several correlated random variables. Therefore, analyzing each of these variables separately cannot provide a complete risk assessment and may result in loss of life or property associated with an underestimation, or an increase in the costs of hydraulic structure associated with an overestimation. Multivariate *FA* (*MFA*) overcomes this problem by simultaneously considering these variables.

Classical *FA* is based on three assumptions namely homogeneity, stationarity and independence. However, different non-standard conditions such as topographic complexity, disturbances by urban development and climate change can influence the hydrological response. Such conditions make flood predetermination, by conventional *MFA* methods, an ineffective exercise and unsuitable to such contexts. The main objective of this thesis is to propose new and more promising methods for the analysis and modeling of hydrological variables both in a multivariate framework and in the lack of the homogeneity assumption. These methods aim to overcome the limitations of classical methods used in the literature. These new proposed methods, based on copulas, allow a better estimation of the risks associated to hydrological extremes by taking into account the dependence between the different variables. As a result, water resource managers can make informed decisions that are better suited to the current context.

In the first part, we propose a novel statistical test for multivariate heterogeneity detection, based on copula and multivariate L-moments. A simulation study is conducted to evaluate the

performance of the proposed test and to compare it with those of existing tests. Results show the ability of the proposed test to discriminate homogeneous and inhomogeneous series. In the second part, we propose a new model based on mixtures of copulas that take into account the heterogeneity of multivariate hydrological series. To estimate the components of this model, we propose a new parameter estimation approach, based on the maximum pseudo-likelihood using genetic algorithms. Results indicate that the proposed method estimates more accurately the parameters even with small sample sizes compared to the existing classical *EM* method. In the last part, we introduce new goodness of fit (*GOF*) tests specifically for multiparameter copulas and adapted to hydrometeorological context. More precisely, the proposed *GOF* tests are based on multivariate L-moments. A simulation study is conducted to evaluate and compare the performances of the proposed tests. The results confirm the usefulness of the new *GOF* tests in comparison with some well-established ones.

Keywords: Frequency analysis; Homogeneity; Dependence; Multiparameter copula; Mixture model; Estimation; Goodness of fit tests; Multivariate L-moments; Flood; Maximum pseudo-likelihood; Meta-Heuristic; Genetic algorithms; Expectation-Maximization.

Avant-Propos

Cette thèse présente les travaux de recherche menés au cours de mes études doctorales. La présente thèse suit la structure standard des thèses par articles de l'INRS-ETE. La première partie de la thèse comporte une synthèse générale qui a pour but de présenter les objectifs et les méthodologies adoptées au cours de la thèse. La deuxième partie de la thèse contient trois articles, dont deux publiés et un autre soumis à des revues internationales avec comité de lecture. Les conclusions ainsi que les perspectives des travaux réalisés sont présentées dans la troisième partie.

Remerciements

Je voudrais adresser mes remerciements les plus sincères à de nombreuses personnes, qui sans leurs soutiens, cette thèse n'aurait pu être accomplie. Je tiens d'abord à remercier mon superviseur Fateh Chebana pour son soutien constant au cours des dernières années. Ses précieux conseils et suggestions ont permis à ces recherches de progresser et d'aboutir. Je désire donc remercier Fateh, mon directeur et mon mentor, pour avoir dirigé ce travail de thèse et pour m'avoir aidée à me préparer pour ma carrière.

J'aimerais également remercier les membres du jury : Jean-François Quessy, Taoufik Bouerzmarni et Sophie Duchesne, pour le temps consacré à l'évaluation du présent travail.

Je tiens aussi à remercier Taha Ouarda, pour ses encouragements, ses conseils et son soutien ainsi que tous les membres du groupe de recherche en hydrologie statistique pour leur coopération et leur aide.

Je remercie mes parents, mes frères, ma sœur et mes beaux-parents pour leur soutien, malgré la distance. C'est à eux que je dédie cette thèse.

Merci à mon mari, Haithem, pour sa patience et sa compréhension, à mes enfants Adam, Amir et Sarah, qui ont toujours été là pour me remonter le moral dans les périodes les plus difficiles et toujours présents pour célébrer mes réussites.

Table des matières

Résumé	ii
Abstract	iv
Avant-Propos	vi
Remerciements	vii
Liste des figures	xiii
Liste des tableaux	xv
Liste des abréviations	xvii
Première partie	1
Chapitre 1 : synthèse	1
1. Introduction et mise en contexte	2
2. Revue de littérature	5
2.1. Introduction aux copules	6
2.2. Les L-moments multivariés.....	10
2.3. Tests de détection des ruptures dans les séries multivariées	11
2.4. Les modèles mixtes multivariés	13
2.5. Tests d'adéquation pour les copules.....	14
3. Problématiques	15
3.1. Détection des points de rupture multivariés en hydrologie.....	15
3.2. Modélisation de l'hétérogénéité dans les séries hydrologiques multivariées	18

3.3. Tests d'adéquation pour les copules multiparamètres.....	21
4. Objectifs et méthodologie	22
4.1. Test de détection des points de rupture multivariés	23
4.1.1. Description de l'objectif.....	23
4.1.2. Méthodologie	24
4.2. Nouvelle méthode d'estimation des paramètres de copules mixtes	26
4.2.1. Description de l'objectif.....	26
4.2.2. Méthodologie	26
4.3. Nouveaux tests d'adéquation pour les copules multiparamètres	27
4.3.1. Description de l'objectif.....	28
4.3.2. Méthodologie	28
Deuxième partie: les articles scientifiques	31
Chapitre 2. Homogeneity testing of multivariate hydrological records, using multivariate copula	
L-moments	33
Abstract	35
1. Introduction	36
2. Theoretical background.....	41
2.1. Multivariate homogeneity testing problem	41
2.2. Multivariate copula L-moments.....	43
3. Multivariate L-moments-based copula homogeneity test	44

4. Simulation study.....	47
4.1. Simulation design.....	48
4.2. Simulation results.....	52
4.2.1. Nominal level evaluation	53
4.2.2. Power evaluation	54
4.2.3. Comparison	56
5. Applications to Real Hydrological Data	60
6. Conclusion.....	65
Chapitre 3. Meta-heuristic estimation method for mixture copula models.....	67
Abstract	69
1. Introduction	70
2. Copula and mixture of copulas.....	75
2.1. Copula model	75
2.2. Mixture copula models.....	76
3. Parameter estimation methods	78
3.1. Expectation-Maximization (<i>EM</i>) method.....	79
3.2. Meta-heuristic Maximum Pseudo-Likelihood method	80
4. Evaluation of estimation methods via simulation	81
4.1. Simulations design	82
4.2. Simulation results.....	85

5. Conclusion.....	96
Chapitre 4. Multivariate L-moment based tests for copula selection, with hydrometeorological applications.....	99
Abstract	101
1. Introduction and literature review	102
2. M-copula in modeling hydrological variables	104
2.1. M-copulas.....	104
2.2. Existing copula <i>GOF</i> tests.....	105
2.3. Multivariate L-moments.....	107
3. The proposed <i>GOF</i> tests for <i>M</i> -copulas	108
3.1. Proposed <i>GOF</i> test based on multivariate copula L-moments.....	109
3.2. Extension of the Shih's <i>GOF</i> test to M-copulas	110
4. Simulation study.....	111
4.1. Simulation design.....	111
4.2. Simulation results.....	112
4.2.1. First type error estimates	112
4.2.2. Power evaluation	114
4.2.3. Comparisons.....	116
5. Illustrative case studies.....	118
5.1. Bivariate Flood Frequency Analysis	118

5.2. Bivariate Drought analysis	124
6. Conclusion.....	129
Troisième Partie	131
Chapitre 5. Conclusions et perspectives.....	131
1. Conclusions générales	133
2. Perspectives	139
Quatrième partie	141
Chapitre 6. Références bibliographiques	141

Liste des figures

Figure 1.1. Différentes étapes d'une AFM et différentes contributions de la thèse.....	23
Figure 2.1. Difference between a) homogeneity and b) change-point testing problem	38
Figure 2.2. Diagram of the simulation study.....	51
Figure 2.3. Percentage of well positioned change detected by different statistics for different scenarios	60
Figure 2.4. Comparison between level curves of the fitted copula before and after the change for a) the Ankang and b) Romaine stations.....	64
Figure 2.5. Level curves of the normal and hüslerReiss copulas for a) same Kendall's τ and b) different Kendall's τ	65
Figure 2.6. Comparison between level curves of the multivariate distribution before and after the change, for a) the Ankang and b) Romaine stations	65
Figure 3.1. Flowchart of the simulations study.....	84
Figure 3.2. Relative errors results of MH-MPL and EM estimators for different sample size n , dependence level $\tau = (0.6, 0.8)$ and $\omega=0.5$: a) weight ratio ω , b) first copula parameter θ_1 and c) second copula parameter θ_2	88
Figure 3.3. Relative errors results between MH-MPL and EM estimators for different kendall's τ , $n=150$ and $\omega=0.5$: a) weight ratio ω , b) first copula parameter θ_1 and c) second copula parameter θ_2	92
Figure 3.4. Relative errors results of MH-MPL and EM estimators for different weight ratio ω , dependence level $\tau = (0.6, 0.8)$ and $n=150$: a) weight ratio ω , b) first copula parameter θ_1 and c) second copula parameter θ_2	96

Figure 4.1. Boxplots showing the power of different *GOF* tests various sample sizes and dependence level $\tau=0.6$ 115

Figure 4.2. Boxplots showing the power of different *GOF* tests various dependence levels and sample size $n=100$ 115

Figure 4.3. Scatter plot of observed data and 10 000 generated pairs from the copula fitted: a) for the Romaine station and b) for the Ankang station121

Figure 4.4. Comparison between the empirical estimate and theoretical Kendall’s function for each fitted copula for (a) Romaine data and (b) Ankang data122

Figure 4.5. Scatter plot of observed data and 10 000 generated pairs from the fitted copula, for modeling dependence structure between precipitation and soil moisture anomalies in northern California, USA: a) from 1948 to 2017 and b) from 1988 to 2017.....126

Figure 4.6. Comparison between the empirical estimate and theoretical Kendall’s function for each fitted copula for modeling the dependence structure between precipitation and soil moisture anomalies in northern California, USA: a) from 1948 to 2017 and b) from 1988 to 2017.....128

Liste des tableaux

Table 2.1. <i>HFA</i> steps in the univariate and multivariate frameworks	37
Table 2.2. Overview of existing tests for change detection in the dependence structure	40
Table 2.3. First type error estimates (%) by the proposed test for different sample sizes and different dependence strengths.....	53
Table 2.4. Power estimates (%) of the proposed statistic T_n	54
Table 2.5. Power comparison between different tests when considering the same copula (here Gumbel) before and after the change	56
Table 2.6. Power comparison between different tests, when considering different copulas before and after the change (here Gumbel and Frank)	58
Table 2.7. Execution time for 100 evaluations.....	59
Table 2.8. General information about the Romaine and Ankang stations	60
Table 2.9. Results of the goodness-of-fit test and model selection criteria, for copula selection ..	61
Table 2.10. Results of the goodness-of-fit test and model selection criteria, for univariate distribution selection	62
Table 2.11. Selected univariate distribution and copula	63
Table 3.1. Overview of estimation method for parameter of mixture models	74
Table 3.2. <i>RRMSE</i> (%) and <i>RBIAS</i> (%) of different estimators, over 1000 replications, for different mixture models and different sample size when dependence level $\tau = (0.2, 0.6)$ and weight ratio $\omega=0.5$	87
Table 3.3. <i>RRMSE</i> (%) and <i>RBIAS</i> (%) of different estimators, over 1000 replications, for different mixture models and different dependence level when sample size $n=150$ and weight ratio $\omega=0.5$	89

Table 3.4. RRMSE (%) and RBIAS (%) of different estimators, over 1000 replications, for different mixture models and different weights ω when sample size $n=150$ and dependence level $\tau = (0.6, 0.8)$	94
Table 4.1. Summary of goodness of fit tests, for hydrological applications.....	104
Table 4.2. M-copulas used in this study (Joe, 2014).....	105
Table 4.3. Summary of the proposed <i>GOF</i> tests for <i>M</i> -copulas	108
Table 4.4. Level of the proposed <i>GOF</i> tests (in %) for various sample sizes and dependence level $\tau=0.6$	113
Table 4.5. Level of the proposed <i>GOF</i> tests (in %) for various dependence levels and sample size $n=100$	114
Table 4.6. Performances of proposed and classical <i>GOF</i> tests (in%) in the case of one parameter copulas for sample size $n=100$ and dependence level $\tau=0.6$	118
Table 4.7. Flood peak and volume results of the goodness-of-fit test and model selection criteria, for univariate distribution selection.....	120
Table 4.8. Results of the application of the <i>GOF</i> tests: Estimated copula parameters by Maximum Pseudo-Likelihood method and corresponding <i>p-value</i> based on 1000 parametric bootstrap samples, for a significance level $\alpha= 5\%$, for Romaine and Ankang data series	122
Table 4.9. Precipitation and soil moisture anomalies results of the goodness-of-fit test and model selection criteria, for univariate distribution selection	125
Table 4.10. <i>P-values</i> for the <i>GOF</i> tests for modeling dependence structure of precipitation and soil moisture anomalies in northern California, USA	127

Liste des abréviations

<i>AF</i>	Analyse Fréquentielle
<i>AFU</i>	Analyse Fréquentielle Univariée
<i>AFM</i>	Analyse Fréquentielle Multivariée
<i>AIC</i>	Akaike Information Criterion
<i>CPD</i>	Change-Point Detection
<i>EM</i>	Expectation-Maximization
<i>FML</i>	Full Maximum Likelihood
<i>GA</i>	Genetic Algorithm
<i>GOF</i>	Goodness Of Fit
<i>HFA</i>	Hydrological Frequency Analysis
<i>IFM</i>	Inference Function for Margins
<i>MH</i>	Meta-Heuristic
<i>MH-MPL</i>	Meta-Heuristic Maximum Pseudo-Likelihood
<i>ML</i>	Maximum Likelihood
<i>MPL</i>	Maximum Pseudo-Likelihood
<i>MSDI</i>	Multivariate Standardized Drought Index
<i>RBIAS</i>	Relative <i>BIAS</i>
<i>RE</i>	Relative error
<i>RRMSE</i>	Relative Root-Mean Square Error
<i>VE</i>	Valeur extrême

Première partie

Chapitre 1 : synthèse

Ce chapitre a pour objectif de faire le lien entre les travaux existants dans la littérature de l'analyse fréquentielle multivariée et les approches proposées dans le cadre de cette thèse. Les principales méthodes classiques considérées pour des fins de comparaison y sont brièvement présentées et discutées. De plus, les principaux éléments méthodologiques y sont résumés afin d'offrir une perspective générale du contenu de cette thèse.

1. Introduction et mise en contexte

Les événements hydrologiques extrêmes, tels que les crues et les étiages, peuvent engendrer des dégâts socio-économiques importants. L'estimation du risque associé à ces événements hydrologiques extrêmes est essentielle pour une gestion intégrée des ressources en eau, pour la protection de la population contre les inondations ou encore pour la conception et construction des ouvrages hydrauliques. Pour prévenir ces événements, il est nécessaire de prédire leurs occurrences et fréquences. Dans ce cadre, l'analyse fréquentielle (*AF*) est une des approches les plus utilisées pour la modélisation et l'estimation adéquate des risques hydrologiques (e.g. Benameur *et al.*, 2017; Franks and Kuczera, 2002; Khaliq *et al.*, 2006). En somme, l'*AF* permet de relier l'amplitude des événements extrêmes à leur fréquence d'occurrence à travers des modèles statistiques (e.g. Hamed and Rao, 1999). La précision et la fiabilité d'une telle analyse sont directement liées à la disponibilité et la qualité des données ainsi qu'aux approches utilisées (e.g. Schumann, 2011).

L'*AF* est composée de quatre étapes principales. En bref, la première étape consiste en une analyse descriptive des données. Cette étape se préoccupe de la qualité de la série de données. La vérification des hypothèses de base à savoir la stationnarité, l'homogénéité et l'indépendance constitue la deuxième étape d'une *AF*. La troisième étape est la modélisation et l'estimation des paramètres du modèle alors que la dernière étape consiste à analyser et évaluer les risques. Pour

plus des détails concernant les différentes étapes d'une *AF*, le lecteur peut se référer au livre de Meylan *et al.* (2012).

L'homogénéité représente l'une des hypothèses fondamentales à vérifier en *AF*. Cette hypothèse postule que l'échantillon provient d'une même population correspondant au même processus physique et que toutes les données ont donc les mêmes propriétés statistiques sous-jacentes (e.g. Shin *et al.*, 2015; Smith *et al.*, 2011). L'étude de l'homogénéité des séries hydrologiques est une tâche incontournable pour l'hydrologue dans la mesure où la réponse qu'elle fournit relève de la caractérisation d'une variabilité climatique et ses effets sur les ressources en eau (e.g. Yan *et al.*, 2017; Yan *et al.*, 2019). Par ailleurs, différentes conditions non standards telles que la complexité topographique, les perturbations par les aménagements urbains et les changements climatiques peuvent influencer la réponse hydrologique et causer une hétérogénéité (e.g. Buishand *et al.*, 2013; Peterson *et al.*, 1998).

La plupart des séries hydrologiques contiennent des hétérogénéités causées par plusieurs facteurs (e.g. Alila and Mtiraoui, 2002; Barth *et al.*, 2017). À titre d'exemple, tel qu'expliqué par Beaulieu *et al.* (2009), le déplacement d'une station de mesure peut introduire une rupture dans la série des données, puisque ce déplacement est souvent accompagné par un changement de l'instrumentation, de l'observateur et l'environnement entourant la station. De plus, la complexité naturelle du processus hydrologique, dérivant par exemple de la topographie des bassins versants, leurs formations géologiques ou également de la variation météorologique, est reconnue comme une source d'hétérogénéité dans les séries hydrologiques (e.g. Smith *et al.*, 2011). La nécessité de bien étudier l'homogénéité des séries de données, avant de commencer une modélisation ou une prévision, est d'autant plus importante dans un contexte de changement climatique (e.g. Yan *et al.*, 2016). Cette propriété, contraignante pour la modélisation hydrologique, doit être prise en compte

dans une démarche de modélisation pour aboutir à des modèles non seulement précis en termes de critères de performance, mais également capable de reproduire la dynamique des processus hydrologiques. En effet, la prise en compte de l'hétérogénéité existante dans les séries hydrologiques est primordiale pour une *AF* complète et réaliste afin d'éviter des résultats et des décisions erronées (e.g. Smith *et al.*, 2011; Villarini *et al.*, 2009; Yan *et al.*, 2016).

Généralement, les événements hydrologiques extrêmes sont caractérisés par plusieurs variables dépendantes, comme le volume, la pointe et la durée pour les crues (e.g. Chebana and Ouarda, 2009; Hao and Singh, 2016). De fait, l'analyse de chacune de ces variables séparément ne peut pas fournir une évaluation complète de la sévérité et des risques associés à un événement. Par conséquent, l'*AF* univariée (*AFU*) peut engendrer éventuellement des pertes matérielles, voire humaines, associées à une sous-estimation ou une augmentation des coûts des ouvrages hydrauliques associés à une surestimation. À ce sujet, plusieurs études ont montré que l'*AF* multivariée (*AFM*) permet de pallier ce problème en considérant conjointement l'impact des différentes variables (e.g. De Michele *et al.*, 2005; De Michele *et al.*, 2013; Grimaldi and Serinaldi, 2006; Karahacane *et al.*, 2020).

La dépendance entre les variables est l'élément clé dans une étude multivariée. Traditionnellement, la modélisation de cette dépendance fait appel aux techniques classiques basées sur la distribution normale multivariée (e.g. Von Storch and Zwiers, 2001; Wilks, 2011). Toutefois, dans le cas des extrêmes hydrologiques, la loi normale n'est pas appropriée et s'ajuste mal aux données (e.g. Zhang and Singh, 2006; Zhang and Singh, 2007a). Par ailleurs, elle ne reproduit pas adéquatement les queues des distributions (e.g. Poulin *et al.*, 2007; Salvadori *et al.*, 2007; Schoelzel and Friederichs, 2008). Pour éviter ce problème, des extensions multivariées de distributions univariées, telles que les lois Gamma et Student, ont été proposées (e.g. Ashkar and Aucoin, 2011; Ma *et al.*, 2013; Yue

et al., 2001). Cet usage des extensions est souvent justifié par la simplicité relative des calculs à effectuer. Cependant, un inconvénient majeur de ces extensions multivariées est que les distributions marginales doivent être identiques et appartenir à la même famille (e.g. Yue *et al.*, 1999b; Yue *et al.*, 2001). De plus, la mesure de dépendance fréquemment utilisée avec ces extensions est le coefficient de corrélation de Pearson. Cependant, l'utilisation de cet indicateur n'est judicieuse que lorsque la relation de dépendance est linéaire et l'univers considéré est gaussien (e.g. Barber *et al.*, 2020; Kelly and Krzysztofowicz, 1997). Ainsi, en hydrologie, domaine où les hypothèses de normalité et de linéarité sont rarement vérifiées, le coefficient de corrélation linéaire n'est pas approprié (e.g. Durocher *et al.*, 2015). Pour répondre à ces attentes, les copules se sont imposées comme un outil incontournable permettant de modéliser la structure de dépendance entre les variables, sans tenir compte des lois marginales (e.g. Genest and Chebana, 2016; Klein *et al.*, 2011; Nelsen, 2013).

La synthèse de cette thèse est organisée comme suit : la section 2 présente une revue de littérature portant sur les approches statistiques utilisées en *AFM* en général. La section 3 résume les différentes problématiques liées aux approches existantes en *AFM*. Les objectifs ainsi que les méthodologies adoptées pour atteindre ces objectifs de la thèse sont présentés dans la section 4.

2. Revue de littérature

Cette section contient la formulation des méthodes d'*AFM* couramment utilisées dans la littérature statistique et/ou hydrologique. Elle est divisée en quatre sous-sections. La première sous-section est consacrée à l'introduction de la notion de copule. La deuxième sous-section présente les méthodes existantes de détection de l'hétérogénéité. La troisième sous-section introduit les modèles mixtes multivariés. Finalement, la dernière sous-section est consacrée à la description des tests d'adéquation pour la sélection des copules.

2.1. Introduction aux copules

La copule est un outil puissant et flexible, permettant l'analyse de la structure de dépendance entre plusieurs variables aléatoires indépendamment des distributions marginales (e.g. Fu and Butler, 2014; Salvadori and De Michele, 2004). Grâce à la notion de copule, nous pouvons donc former des distributions multidimensionnelles avec différentes distributions marginales. Formellement, la copule est une fonction de répartition à plusieurs dimensions qui exploite la dépendance entre plusieurs variables afin d'obtenir une distribution de probabilité multivariée à partir de leurs distributions marginales (e.g. Grimaldi and Serinaldi, 2006; Klein *et al.*, 2011). Par cet intermédiaire, il est possible de produire des algorithmes de simulations d'évènements hydrologiques se rapprochant davantage de la réalité multidimensionnelle du phénomène comme les crues (e.g. Vezzoli *et al.*, 2017; Vittal *et al.*, 2015). Pour une revue générale des familles de copules, le lecteur peut se référer aux ouvrages de Joe (2014) et Nelsen (2013).

Selon le théorème de Sklar (1959), que nous présentons ici par souci de clarté dans le cadre bivarié, pour un vecteur de variables aléatoires (X_1, X_2) ayant comme fonctions de répartitions marginales (F_1, F_2) et comme fonction de répartition jointe F , il existe une copule C qui vérifie

$$F(x_1, x_2) = C(F_1(x_1), F_2(x_2)) \quad (1.1)$$

Si les marginales (F_1, F_2) sont continues, alors C est unique. Comme la copule C est invariante par transformation croissante des marges, en effectuant le changement de variables $u_i = F_i^{-1}(x_i)$ pour $i=1,2$, l'équation (1.1) nous donne

$$C(u_1, u_2) = F(F_1^{-1}(u_1), F_2^{-1}(u_2)) \quad (1.2)$$

Une conséquence du théorème de Sklar est que plusieurs mesures de dépendance appropriée pourraient s'exprimer uniquement en fonction de la copule C , car aucune information sur la

dépendance n'est incluse dans les marges (e.g. Genest and Rivest, 1993). Les coefficients les plus utilisés pour quantifier la dépendance dans les séries multivariées sont le tau de Kendall et le rho de Spearman.

Il existe plusieurs types de copules ayant des propriétés et applications différentes, dont les copules archimédiennes et valeurs extrêmes. Pour plus de détails concernant ces deux classes de copules, le lecteur peut se référer aux livres de Nelsen (2013) et Salvadori *et al.* (2007). En hydrologie, les copules archimédiennes sont couramment utilisées grâce à la simplicité de leur forme et de leur flexibilité d'application. De plus, elles permettent de construire une grande variété de familles de copules, et donc de représenter une grande variété de structures de dépendance (e.g. AghaKouchak, 2014; Zhang and Singh, 2007a). En effet, contrairement aux copules gaussiennes, les copules archimédiennes ont le grand avantage de décrire des structures de dépendance très diverses dont notamment les dépendances dites asymétriques, où les coefficients de queue inférieure et supérieure sont différents. Dans cette classe de copule, il y a des copules à un seul paramètre tels que Clayton, Frank et Gumbel, et d'autres multiparamètres tels que *BB1* (Clayton-Gumbel), *BB6* (Joe-Gumbel) et *BB7* (Joe-Clayton). Les copules multiparamètres sont de plus en plus utilisées en hydrologie grâce à leur flexibilité à modéliser simultanément les dépendances aux queues inférieures et supérieures (e.g. De Michele *et al.*, 2013; Requena *et al.*, 2016; Sadegh *et al.*, 2017; Salvadori and De Michele, 2010).

La copule valeur extrême (*VE*) permet de mesurer la dépendance des événements rares, ce qui est essentiel par exemple pour l'étude de la crue. Ce type de copules est défini à partir de la fonction de dépendance de Pickands $A(\cdot)$ (e.g. Gudendorf and Segers, 2010; Salvadori and De Michele, 2010). Les propriétés de la fonction de dépendance A sont données dans Pickands (1989) et Capéraà *et al.* (1997). L'avantage d'utilisation de cette famille de copules réside dans le fait que pour

construire une distribution de valeurs extrêmes bidimensionnelle, il suffit ainsi de coupler des marges issues des lois de la théorie de valeurs extrêmes avec une copule *VE* (e.g. Genest and Nešlehová, 2014). Parmi les modèles de dépendance de nature extrême, il existe des copules à un seul paramètre, telles que les copules de Gumbel et Galambos, ainsi que des copules multiparamètres tel que la copule *BB5* (e.g. Joe, 2014).

L'estimation des paramètres d'une copule se présente comme une étape importante dans le processus de modélisation. Dans la littérature, il existe plusieurs approches d'estimation qui peuvent être classées en approches paramétriques et semi-paramétriques. La méthode du maximum de vraisemblance complète (Full Maximum Likelihood, *FML*), proposée par Shih and Louis (1995), est l'une des approches paramétriques les plus utilisées (e.g. Kim *et al.*, 2007). Elle consiste à estimer conjointement les paramètres des distributions marginales et ceux de la copule. L'inconvénient de cette méthode est que la vraisemblance peut être difficile à maximiser et le problème d'optimisation est difficile à résoudre. De plus, pour les copules multiparamètres, la méthode nécessite des temps de calcul très longs et intensifs (e.g. Zhang and Singh, 2019).

Afin de remédier à ces inconvénients, la méthode d'inférence de marges (Inference Function for Margins, *IFM*) a été développée par Joe and Xu (1996). Cette méthode repose sur le fait que la représentation en copule permet de séparer les paramètres des distributions marginales de ceux de la copule. Ainsi, elle estime dans un premier temps les paramètres des marges et ensuite ceux de la copule, en utilisant dans les deux étapes la méthode du maximum de vraisemblance. Toutefois, tout comme la méthode *FML*, la méthode *IFM* souffre des inconvénients des approches paramétriques. En effet, les approches paramétriques dépendent des hypothèses de distributions marginales et sont sensibles aux spécifications de ces dernières. Particulièrement, puisque les marges interviennent directement dans le calcul de la vraisemblance, une éventuelle erreur d'estimation des fonctions

marginales peut se propager et rendre ainsi erronée l'estimation des paramètres de la copule (e.g. Joe and Xu, 1996).

Tel que discuté par Kim *et al.* (2007), si les modèles paramétriques de la copule et/ou des marges sont mal spécifiés, les approches paramétriques conduisent généralement à un mauvais ajustement aux données. Dans ce cas, les méthodes d'estimation semi-paramétriques peuvent être une bonne alternative. En effet, l'approche semi-paramétrique ne fait pas d'hypothèses sur les distributions marginales. Elle se base sur une estimation empirique de ces dernières. Pour cela, les observations sont transformées en données uniformes en utilisant les fonctions empiriques basées sur les rangs des observations. Les méthodes semi-paramétriques les plus connues sont la méthode des moments et la méthode de maximum pseudo-vraisemblance (e.g. Brahimi and Necir, 2012; Genest and Chebana, 2017).

La méthode des moments est basée sur l'inversion d'un coefficient de dépendance tel que le tau de Kendall ou le rho de Spearman. Cette méthode utilise la relation, plus ou moins explicite, qui peut exister entre le coefficient de dépendance et le paramètre de la copule. Pour ce faire, la version empirique du coefficient de dépendance, basé sur le rang des observations, est égalisée à la version théorique de ce dernier (e.g. Genest *et al.*, 2011b; Kojadinovic and Yan, 2010). L'estimateur obtenu est consistant et asymptotiquement non biaisé et normal. Cependant, cette méthode ne s'applique que lorsqu'il n'y a qu'un seul paramètre à estimer.

La méthode du maximum de pseudo-vraisemblance (Maximum Pseudo-Likelihood, *MPL*) permet de combler les limitations de la méthode des moments (e.g. Hoff, 2007). En effet, elle permet d'estimer les paramètres lorsqu'il s'agit d'une copule multiparamètre. Elle est également basée sur les rangs des observations. Ainsi, elle est robuste contre les erreurs de spécification des distributions marginales. Récemment, une nouvelle méthode basée sur les L-moments multivariés

a été développée par Brahim *et al.* (2015), spécifiquement pour l'estimation des copules multiparamètres. L'avantage de cette méthode est que l'estimateur est non biaisé et robuste à l'existence des points aberrants. En hydrologie, l'utilisation des L-moments multivariés est un avantage important dans la modélisation des distributions à queue lourde (Serfling and Xiao, 2007). La bonne performance des méthodes basées sur les L-moments, dans le cas des séries de faible taille, rend les L-moments de plus en plus utiles en hydrologie.

2.2. Les L-moments multivariés

Les L-moments sont des combinaisons linéaires des moments de probabilité pondérés. La théorie des L-moments a été développée par Hosking (1990) pour résoudre certains problèmes liés à l'ajustement des lois statistiques. L'extension multivariée des L-moments a été proposée par Serfling and Xiao (2007). Parmi les avantages des L-moments, on note qu'ils existent si et seulement si la moyenne de la distribution existe. De plus, la distribution est bien caractérisée par ses L-moments. Par ailleurs, les L-moments sont utiles pour caractériser les distributions multivariées à queue lourde et dans le cas des petites tailles d'échantillons. Notons aussi que les estimateurs de L-moments possèdent l'avantage d'être non biaisés (Brahimi *et al.*, 2015; Serfling and Xiao, 2007). Dans ce qui suit, nous procédons à une présentation sommaire des L-moments dans un cadre bivarié. On pourra éventuellement se référer à Serfling and Xiao (2007) pour une présentation plus détaillée.

Soit $X^{(j)}$ une variable aléatoire ayant comme fonction de distribution F_j , pour $j=1,2$. Les L-moments multivariés d'ordre k sont les matrices λ_k dont les éléments sont définis par :

$$\lambda_{k[ij]} = Cov(F_i(X^{(i)}), P_{k-1}^*(F_j(X_j))), i, j = 1,2 \text{ et } k = 2,3, \dots \quad (1.3)$$

tel que P_{k-1}^* sont les polynômes de Légendre décalés (Chang and Wang, 1983). En particulier, les trois premiers L-moments bivariés sont définis par :

$$\begin{aligned}\lambda_{2[12]} &= 2Cov(F_1(X^{(1)}), F_2(X^{(2)})) \\ \lambda_{3[12]} &= -6Cov(F_1(X^{(1)}), F_2(X^{(2)})(1 - F_2(X^{(2)}))) \\ \lambda_{4[12]} &= Cov(F_1(X^{(1)}), 20F_2^3(X^{(2)}) - 30F_2^2(X^{(2)}) + 12F_2(X^{(2)}) - 1)\end{aligned}\quad (1.4)$$

Tel que établi par Brahim *et al.* (2015), les L-moments bivariés d'ordre k peuvent être définis en fonction de la copule C par :

$$\lambda_{k[12]}^C = \int_0^1 \int_0^1 (C(u_1, u_2) - u_1 u_2) du_1 dP_k^*(u_2) \quad (1.5)$$

Afin d'avoir une indépendance des L-moments d'ordres supérieurs ($k \geq 3$) par rapport aux unités de mesure des variables $X^{(1)}$ et $X^{(2)}$, leur forme standardisée est utilisée. On définit alors les rapports de L-moments par :

$$\tau_{k[12]}^C = \frac{\lambda_{k[12]}^C}{\lambda_{2[12]}^C}, \text{ for } k \geq 3 \text{ and } \tau_{2[12]}^C = \frac{\lambda_{2[12]}^C}{\lambda_{1[2]}^C} \quad (1.6)$$

Le rapport τ_3 est une mesure de l'asymétrie et τ_4 est une mesure de l'aplatissement.

2.3. Tests de détection des ruptures dans les séries multivariées

Dans l'AF classique, l'homogénéité est l'une des hypothèses importantes à respecter. Cependant, cette hypothèse n'est pas toujours réaliste. À titre indicatif, si l'évènement extrême provient du mélange de deux phénomènes physiques générateurs différents, par exemple crues pluviales ou fonte nivale, la série peut présenter une hétérogénéité significative (e.g. Peterson *et al.*, 1998).

Des techniques statistiques peuvent être utilisées pour détecter les hétérogénéités, en utilisant la théorie classique de détection des ruptures. La détection des ruptures consiste à déceler la présence d'un ou plusieurs changements brutaux dans la loi des observations et à les localiser dans la série (e.g. Lung-Yut-Fong *et al.*, 2011; Rougé *et al.*, 2013). Si ces changements ne sont pas détectés et pris en considération lors de l'analyse, les résultats peuvent être biaisés et la décision qui en découle sera erronée (e.g. Ehsanzadeh *et al.*, 2011). Dans la littérature, une panoplie de tests paramétriques et non paramétriques est proposée pour la détection des points de rupture dans une série (e.g. Bouzebda and Keziou, 2013; Kundzewicz and Robson, 2004; Li and Liu, 2004). Les tests favorisés en AF en général et dans cette étude en particulier sont de type non paramétrique. En effet, les tests paramétriques impliquent beaucoup de choix arbitraires, par exemple le choix du noyau, de la fonction de pondération ou encore la dimension de la fenêtre de calcul, ce qui rend leur utilisation trop contraignante et inadéquate (e.g. Xie *et al.*, 2014). De plus, les approches non paramétriques évitent l'imposition d'une distribution paramétrique aux données et les hypothèses sous-jacentes. Ceci est avantageux puisqu'en AF l'étape de vérification de l'homogénéité précède l'étape de modélisation et le choix des distributions.

Dans le cadre d' AFU , plusieurs tests non paramétriques pour la détection des points de rupture ont été développés en statistique et appliqués en hydrologie. Pour une revue de littérature de ces différentes techniques, ainsi que des études de puissance, le lecteur peut se référer à Bryden *et al.* (1995). Toutefois, la littérature concernant la détection de points de rupture dans le cadre multivarié est moins étoffée que dans le cas univarié. Néanmoins, quelques tests multivariés de détection des points de rupture ont été développés. Parmi ces tests, il y a les tests basés sur la notion de profondeur statistique. Chebana *et al.* (2017) a fait une revue de littérature et une comparaison de ces différents tests avec des simulations en plus des applications hydrologiques. Une autre catégorie de tests,

basée sur l'extension du test univarié de Pettitt (1979) au contexte multivarié, est également proposée pour la détection des ruptures dans les séries multivariées (e.g. Holmes *et al.*, 2013). Dias and Embrechts (2004) ont également proposé un test basé sur le rapport de vraisemblance. Plus récemment, Quessy *et al.* (2013) et Kojadinovic *et al.* (2016) ont proposé des tests de détection de la rupture dans la structure de dépendance basée sur les processus de Kendall et Spearman, respectivement.

2.4. Les modèles mixtes multivariés

Si une hétérogénéité est détectée dans la série des données, il est nécessaire de l'incorporer dans la modélisation. Dans un contexte univarié, la modélisation tenant compte de l'hétérogénéité est largement étudiée avec diverses approches. La première approche consiste à séparer la série des données selon les processus saisonniers (par exemple crue printanière ou automnale) (e.g. Villarini and Smith, 2010). Or, cette approche peut ne pas être adaptée au contexte hydrologique. Par exemple, dans le cas de la méthode de dépassements de seuil, cette approche va réduire davantage la taille de la série. La deuxième approche est basée sur l'utilisation des modèles mixtes. Dans ce contexte, plusieurs modèles de distributions mixtes sont développés et appliqués en hydrologie (e.g. Alila and Mtiraoui, 2002; Yan *et al.*, 2016).

Dans un contexte univarié, plusieurs études ont mis l'accent sur l'importance de considérer des distributions mixtes pour la modélisation de l'hétérogénéité en *AF* (e.g. Singh *et al.*, 2005a; Villarini and Smith, 2010). Toutefois, peu de modèles ont été construits pour tenir compte de l'hétérogénéité dans un contexte multivarié en général et particulièrement en hydrologie. À cet égard, quelques études ont suggéré l'utilisation des copules mixtes (e.g. Thongkairat *et al.*, 2019; Vrac *et al.*, 2012; Yu *et al.*, 2013). En pratique, lors de la mise en œuvre de tels modèles, l'un des principaux défis est l'estimation de ses paramètres. Pour avoir un meilleur ajustement aux données,

il incombe d'avoir une méthode pour estimer efficacement les différents paramètres du modèle. Dans la littérature, la méthode du maximum de vraisemblance, combinée avec l'algorithme de maximisation de l'espérance (EM), est la méthode la plus utilisée pour estimer les paramètres des copules mixtes (e.g. Ding and Song, 2016). Il s'agit d'une méthode itérative en deux étapes : le calcul de l'espérance (E) et sa maximisation (M). Le principe de l'algorithme est de maximiser de manière itérative l'espérance du log-vraisemblance conditionnellement aux valeurs des paramètres calculées à la première étape. L'algorithme alterne entre ces deux étapes jusqu'à la convergence (c'est-à-dire la variation de la vraisemblance est inférieure à un seuil prédéfini).

2.5. Tests d'adéquation pour les copules

En choisissant une famille paramétrique pour la copule d'intérêt, il est naturel de se demander si cette copule explique bien la dépendance entre des variables étudiées. Pour répondre à cette question, en général, le choix de la meilleure copule se fait à l'aide des tests d'adéquation (e.g. Berg, 2009; Fermanian, 2005; Fermanian, 2013). Ces derniers permettent de vérifier formellement si la copule sous-jacente à une population appartient ou non à une certaine famille paramétrique de copules $\mathfrak{C} = \{C_\theta, \theta \in \Theta\}$ où $\Theta \subset \mathbb{R}^p$ est l'espace des paramètres. Ainsi, l'hypothèse nulle $H_0: C \in \mathfrak{C}_\theta$ est testée contre l'alternative $H_1: C \notin \mathfrak{C}_\theta$.

Différents tests d'adéquation pour les copules sont proposés dans la littérature. Le développement de ces tests est relativement récent et toujours en plein essor (e.g. Berg and Quessy, 2009; Durocher and Quessy, 2017). Parmi les tests les plus utilisés, nous trouvons ceux basés sur le processus de la copule empirique, développés par Genest *et al.* (2009). Ces tests consistent à calculer une distance entre la copule empirique et une estimation de la copule paramétrique obtenue sous l'hypothèse nulle. D'une manière analogue, d'autres tests basés sur le processus de Kendall, appelé aussi la transformation intégrale de probabilité, ont été proposés (e.g. Genest *et al.*, 2006). Pour ces

tests, plusieurs versions de statistique peuvent être construites, dont les plus connues sont celles de Kolmogorov-Smirnov et de Cramer-Von-Mises.

3. Problématiques

La problématique générale de la thèse repose sur le constat que les approches disponibles dans la littérature et couramment utilisées en *AFM* sont inappropriées pour diverses situations réalistes telles que la faible taille d'échantillon, l'existence des ruptures dans les séries de données ainsi que l'hétérogénéité des séries hydrologiques.

Dans un cadre univarié, les différentes étapes d'une *AF* sont largement étudiées. Pour plus de détails, le lecteur peut se référer, à titre d'exemple, à l'ouvrage de Rao and Hamed (2000). Cependant, ces étapes sont moins considérées dans un contexte multivarié. Certes, durant les dernières années, le développement de l'étape de la modélisation, dans un contexte multivarié, a reçu un intérêt grandissant. À ce sujet, de plus en plus d'études sont apparues (e.g. Fan *et al.*, 2015; Karahacane *et al.*, 2020; Requena *et al.*, 2013b; Vittal *et al.*, 2015). Toutefois, comme souligné dans Chebana (2013), les étapes précédant la modélisation multivariée, tout particulièrement la vérification des hypothèses, sont négligées, et ceci malgré leur importance. La contribution de la présente thèse est de combler ce vide par la détermination, l'adaptation et l'application des techniques multivariées pour certains de ces étapes et éléments manquants pour une *AFM* complète. Dans ce qui suit, nous allons aborder plus spécifiquement trois sous-problématiques à l'intérieur de cette problématique générale.

3.1. Détection des points de rupture multivariés en hydrologie

Pour que les résultats d'une *AFM* soient valables, la série de données doit satisfaire un certain nombre de conditions, dont l'homogénéité. Or, en hydrologie, les séries de maximums annuels sont

parfois composées d'évènements issus de différents processus (e.g. Yan *et al.*, 2019). Par conséquent, les crues provoquées par ces différents évènements peuvent avoir des distributions très différentes (e.g. Barth *et al.*, 2017; Yan *et al.*, 2017). À cet égard, afin de détecter les changements dans les séries multivariées, plusieurs tests ont été proposés. Parmi ces tests, il y a ceux basés sur les fonctions de profondeur, tels que le test M et le test T, développés par Li and Liu (2004), et les tests d'indice de qualité, développés par Liu and Singh (1993). Pour plus de détails sur ces tests ainsi qu'une revue bibliographique sur les études de comparaisons de performance entre ces tests, le lecteur peut se référer à Chebana *et al.* (2016). Toutefois, ces tests ont des limitations et des contraintes. En effet, ils permettent de détecter seulement les ruptures au niveau des paramètres de la distribution. Or, dans certains cas, le changement peut affecter la distribution elle-même et non seulement ses paramètres (e.g. Vezzoli *et al.*, 2017). En plus, ces tests ne permettent pas de déceler le type de rupture c'est-à-dire la différenciation entre les ruptures au niveau des fonctions marginales et celle de la fonction de dépendance. Par conséquent, aucun renseignement sur les possibles causes ayant initié ces ruptures n'est obtenu.

Pour remédier à ce problème, des efforts ont été déployés dans des études antérieures afin de proposer des tests de détection de la rupture dans la structure de dépendance. Une première classe de tests, basée sur le rapport de vraisemblance, a été développée (e.g. Bouzebda and Keziou, 2013; Dias and Embrechts, 2004; Guegan and Zhang, 2010). Néanmoins, ces tests sont critiquables à plusieurs niveaux. En effet, certaines versions de ces tests sont de type paramétrique ou semi-paramétrique au sens où ils supposent que les fonctions marginales sont connues et ayant une forme précise, ce qui n'est pas toujours réaliste d'un point de vue pratique. En outre, certains tests supposent que la copule est *a priori* connue, ce qui n'est pas toujours le cas. De plus, les séries de données hydrologiques ne satisfont pas à certaines conditions nécessaires à la stricte applicabilité

de ces tests, tel que par exemple la grande taille d'échantillon et la forte dépendance entre les variables. Pour cette raison, Holmes *et al.* (2013) suggèrent d'utiliser des tests de type non paramétrique. Cependant, un inconvénient majeur de ces tests est que les marges sont supposées homogènes. De ce fait, nous ne pourrions conclure en faveur d'une rupture dans la copule qu'en faisant l'hypothèse supplémentaire qu'il n'y a pas de rupture dans les marges des observations. Donc, si une des séries univariées n'est pas homogène, cela pourrait biaiser la détection de rupture dans la dépendance.

Dans la même optique, des tests non paramétriques de détection de rupture, particulièrement sensibles à un changement dans le rho de Spearman ou le tau de Kendall, sont proposés par Kojadinovic *et al.* (2016) et Quessy *et al.* (2013), respectivement. Un problème qui survient avec ces tests réside dans le fait que, par construction, ils n'auront aucune puissance pour détecter les ruptures si la valeur de rho de Spearman ou de tau de Kendal demeure constante. De plus, des simulations, visant l'étude de performance de ces tests, ont révélé que ces tests semblent peu puissants pour les échantillons de taille modérée à faible (e.g. Dehling *et al.*, 2017; Wied *et al.*, 2014). Or, des séries de données larges ne sont toutefois pas la norme en *AF* hydrologique, ce qui rend ces tests moins adaptés pour les applications hydrologiques. Subséquemment, l'existence de plusieurs points de rupture dans la structure de dépendance est envisageable, particulièrement sur de longues séries de données. Dans ce cas, ces tests sont peu ou pas adaptés. En effet, afin de détecter des ruptures multiples, ces tests sont souvent appliqués en utilisant une procédure de segmentation. Par conséquent, leur puissance diminuera puisqu'ils sont utilisés sur des échantillons de plus en plus petits.

En résumé, les postulats sur lesquels reposent la plupart des tests cités précédemment ne sont pas toujours vérifiés dans le contexte hydrologique. Rappelons que, même si ces tests peuvent être

appliqués à un problème de détection des ruptures multivariées, ils souffrent d'un inconvénient majeur du fait qu'ils sont tous basés sur l'hypothèse que le changement affecte seulement le paramètre de la copule. Le type de la copule est supposé constant, ce qui n'est pas toujours justifiable en pratique. Par conséquent, des procédures plus flexibles, capables de détecter les ruptures dans la fonction de dépendance, seraient avantageuses.

3.2. Modélisation de l'hétérogénéité dans les séries hydrologiques multivariées

L'homogénéité des séries d'observations est l'une des hypothèses fondamentales dans une *AF* (e.g. Rao and Hamed, 2000). Or, au cours des dernières décennies, cette hypothèse est désormais reconnue comme étant inadéquate en raison des changements naturels dus aux événements issus de processus différents ou des changements anthropiques dus aux activités humaines (e.g. Alila and Mtiraoui, 2002; Beaulieu *et al.*, 2009). La seule détection des ruptures au sein des séries multivariées n'est pas suffisante. Il est nécessaire d'inclure cette hétérogénéité dans la modélisation afin d'établir des stratégies adéquates de gestion des risques associés aux événements extrêmes (e.g. Evin *et al.*, 2011; Shin *et al.*, 2015; Smith *et al.*, 2011; Yan *et al.*, 2016). Cependant, dans un contexte multivarié, seulement quelques études portent sur le sujet, en dehors du domaine de l'hydrologie (e.g. Vrac *et al.*, 2012; Vrac *et al.*, 2005; Yu *et al.*, 2013).

Dans la littérature, c'est surtout en finance que l'aspect hétérogène des données est intégré dans l'étape de modélisation (e.g. Guegan and Zhang, 2010; Hu, 2006). En ce sens, quelques modèles de copules mixtes sont développés pour les séries non homogènes. Généralement, ces modèles sont basés sur l'hypothèse de normalité. Or, une telle hypothèse pourrait être inadéquate dans certains cas et particulièrement dans le cas des variables hydrologiques. Au vu de ce constat, d'autres modèles de copules mixtes ont été proposés (e.g. Christensen *et al.*, 2019; Qu and Lu, 2019; Vrac *et al.*, 2005). Un inconvénient majeur de ces modèles est que la structure de dépendance est

supposée constante, c'est-à-dire que les deux copules, avant et après le point de rupture, appartiennent à la même famille. Toutefois, en hydrologie, cette hypothèse pourrait poser un problème. En effet, les séries de maximums annuels sont souvent composées d'évènements issus de processus tout à fait différents. Par conséquent, les crues provoquées par ces différents évènements peuvent avoir des distributions très différentes. Il est à noter que depuis peu, quelques études se sont intéressées à la modélisation mixte des séries hydrologiques multivariées (e.g. Fan *et al.* (2016), Khan *et al.* (2019) et Li *et al.* (2013)). Toutefois, il importe de mentionner que ces études considèrent seulement les marges comme étant hétérogènes en faisant l'hypothèse que la structure de dépendance est homogène. Dans ce sens, le modèle consiste à utiliser une combinaison des distributions univariées mixtes pour les marges et une seule copule pour la dépendance.

Dans cette lignée, plusieurs auteurs ont proposé de partitionner les séries de données afin d'obtenir des sous-séries homogènes. Récemment, quelques modèles de copules saisonnières ont été développés pour les séries multivariées hétérogènes représentant une saisonnalité (e.g. Christensen *et al.*, 2019; Gómez *et al.*, 2017). En revanche, cette approche pourrait ne pas être adaptée pour l'estimation des évènements extrêmes ou ne pas conduire nécessairement aux meilleurs résultats en traitant les processus hydrologiques (e.g. Barth *et al.*, 2017). En effet, les séries disponibles sont plutôt courtes, inférieures à une centaine d'années pour la grande majorité d'entre elles. Or, les courtes séries sont inappropriées pour obtenir des estimations fiables des quantiles, en particulier extrêmes. Ce problème persiste et s'aggrave lors de l'élaboration d'un modèle basé sur les séries partielles. En outre, le fait de considérer une série partielle constitue une importante perte d'information et occasionne également une grande incertitude sur les prédictions (e.g. Yan *et al.*, 2019). En somme, le développement de nouveaux modèles dans un cadre d'*AFM* hétérogène répond donc à un besoin qui semble de plus en plus important. En particulier, l'utilisation d'un

modèle de copules mixtes permettrait une exploitation plus complète de l'information disponible, améliorant ainsi les risques associés aux événements extrêmes. Par ailleurs, ce modèle permettrait de modéliser une grande variété de données et d'obtenir une flexibilité additionnelle en comparaison avec les modèles de copules classiques.

La mise en œuvre d'un modèle de copules mixtes passe par l'estimation de ces différents paramètres. Estimer de manière adéquate les paramètres du modèle est l'un des plus importants problèmes de l'inférence statistique en général et en *AF* en particulier. L'estimateur le plus fréquemment utilisé en statistique, et particulièrement dans les applications des modèles mixtes univariés et multivariés, est généralement obtenu par la méthode du maximum de vraisemblance (Maximum Likelihood, *ML*) (e.g. Caudill and Acharya, 1998; Fu *et al.*, 2019; Shin *et al.*, 2014). Pour résoudre de manière efficiente le système du *ML*, l'algorithme itératif *EM* a été largement utilisé (e.g. Arcidiacono and Jones, 2003; Celeux *et al.*, 2001). Rappelons que cet algorithme n'a jamais été considéré dans le cadre d'une *AFM* en hydrologie. L'avantage de l'utilisation de cet algorithme est qu'il garantit que la vraisemblance augmente à chaque itération, ce qui conduit donc à un estimateur de plus en plus précis, lorsque celui-ci existe. Néanmoins, cet algorithme comporte quelques inconvénients. Dans certains cas, l'algorithme peut ne converger que vers un maximum local de la vraisemblance (e.g. Ding and Song, 2016; Dou *et al.*, 2016). Dès lors, l'algorithme *EM* introduit un biais dans l'estimation des paramètres et augmente donc les incertitudes d'estimation des différents quantiles. De plus, l'algorithme dépend des valeurs initiales des paramètres choisis arbitrairement. De ce fait, pour certaines mauvaises valeurs, l'algorithme peut rester gelé en un point, alors qu'il convergera vers le maximum global pour d'autres valeurs initiales plus pertinentes. Par conséquent, l'algorithme *EM* peut parfois nécessiter plusieurs initialisations différentes afin d'assurer sa convergence. Enfin, pour des échantillons de taille restreinte, ce qui

est souvent le cas en hydrologie, l'algorithme est d'une utilité limitée. Par exemple, dans le contexte d'une *AFU*, Shin *et al.* (2014) ont montré que l'estimation par la méthode *ML* couplée avec l'algorithme *EM* est moins performante lorsque la taille de l'échantillon est inférieure à cent et que l'algorithme ne converge pas lorsque la taille de l'échantillon devient plus petite que cinquante. D'où la nécessité de proposer de nouvelles méthodes d'estimation adaptées au contexte hydrologique.

3.3. Tests d'adéquation pour les copules multiparamètres

Pour être en mesure d'estimer adéquatement les quantiles associés aux événements extrêmes, il est nécessaire de choisir le type de copules qui reproduisent le mieux le comportement des variables hydrologiques. Dans la littérature, il existe plusieurs tests statistiques pour déterminer si un échantillon donné est tiré d'une certaine famille de copules. Ces tests permettent de déterminer quelles sont les copules les plus appropriées à la modélisation d'un phénomène donné, en se basant de manière générale sur la comparaison d'un estimateur paramétrique du modèle avec un estimateur non paramétrique (e.g. Fermanian, 2013; Genest *et al.*, 2009; Mesfioui *et al.*, 2009).

Comme mentionné dans la section 2.3, les tests d'adéquation les plus utilisés sont basés sur les processus de copule empirique ou le processus de Kendall (e.g. Berg and Quessy, 2009; Kojadinovic *et al.*, 2011). Toutefois, l'utilisation de ces tests soulève quelques problèmes. En effet, le test basé sur la transformation intégrale de probabilité peut ne pas converger pour quelques types de copules. De plus, toutes les copules de valeurs extrêmes possèdent la même fonction du Kendall (e.g. Gudendorf and Segers, 2010; Salvadori and De Michele, 2010). De ce fait, le test basé sur la transformation intégrale de probabilité n'est pas capable de discriminer entre les copules appartenant à cette famille. Par ailleurs, pour certaines familles de copules, la transformation intégrale de probabilité n'admet pas de forme explicite, ce qui rend difficile l'utilisation de ce test.

De plus, puisque la transformation intégrale de probabilité est une projection unidimensionnelle de la fonction de dépendance, l'efficacité du test décroît pour les copules multiparamètres. Quant au test basé sur le processus de copule empirique, tel que montré dans l'étude de simulations par Genest *et al.* (2009), il n'est pas puissant pour les séries aussi courtes que les séries hydrologiques. Aussi, la performance de ces tests a été évaluée, par simulations, seulement dans le cas des copules à un seul paramètre. Le cas des copules multiparamètres n'a pas été considéré.

4. Objectifs et méthodologie

Étant donné les problèmes liés aux méthodes existantes que l'on peut retrouver en procédant à une *AFM* des variables hydrologiques, l'objectif global de la présente thèse est de proposer des méthodologies plus adaptées au contexte hydrologique qui permettent de contourner les lacunes et limitations des méthodes existantes dans la littérature. Ainsi, on se sert de certaines approches statistiques prometteuses pour pallier les limitations des méthodes existantes et couramment utilisées en *AFM* afin de s'adapter aux spécificités des processus hydrologiques. Bien que les méthodologies proposées dans la présente thèse soient appliquées aux variables hydrologiques, elles se veulent applicables à n'importe quel autre domaine tel que la finance ou l'actuariat. L'objectif global est scindé en trois objectifs spécifiques, lesquels sont présentés de manière détaillée sous forme d'articles dans les chapitres subséquents, donnant lieu à trois contributions. Pour chacune des contributions, une description succincte de la méthodologie utilisée est présentée. La figure 1.1 présente une vue d'ensemble des approches développées et de leur application dans le cadre de cette thèse.

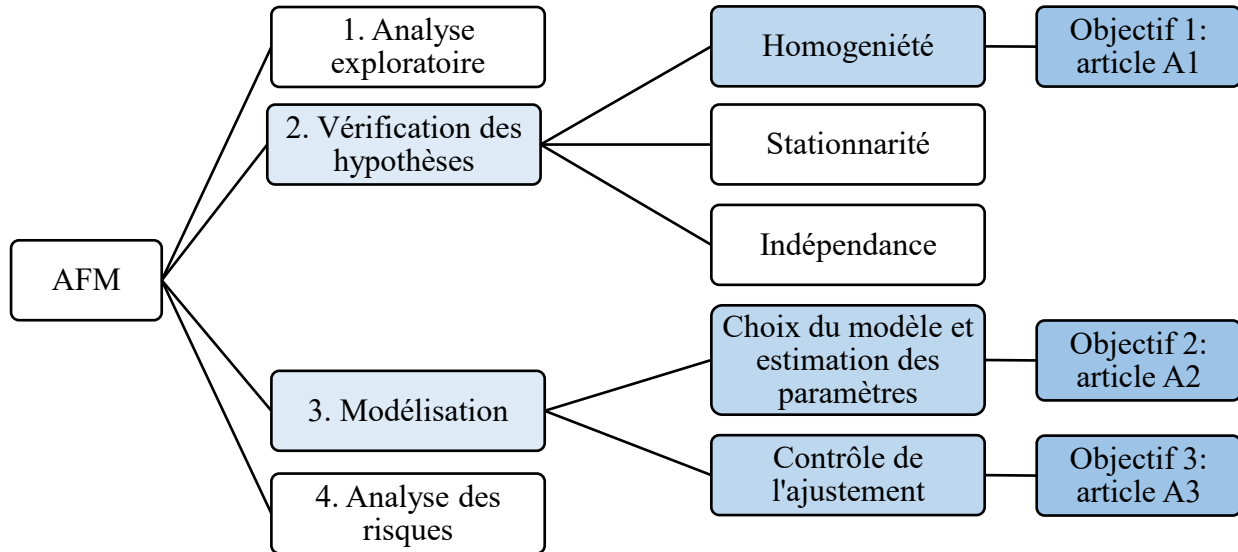


Figure 1.1. Différentes étapes d'une AFM et différentes contributions de la thèse

4.1. Test de détection des points de rupture multivariés

Dans cette section, nous allons présenter l'objectif de l'article [A1] intitulé « *Homogeneity testing of multivariate hydrological records, using multivariate copula L-moments* » ainsi qu'une courte description de la méthodologie. Davantage de détails concernant cette méthodologie sont présentés dans le chapitre 2, correspondant à l'article.

4.1.1. Description de l'objectif

Le premier objectif spécifique est de développer un nouveau test d'homogénéité pour les séries multivariées, visant à combler les lacunes citées plus haut (Section 3.1). Le nouveau test permettra notamment la détection des points de rupture dans la structure de dépendance tout en relaxant les hypothèses des tests existants. Pour ce faire, la structure de dépendance est décrite par la copule, ce qui permet de séparer les liens d'interdépendance des effets marginaux. De plus, le test proposé

est entièrement non paramétrique au sens où aucune copule n'est préalablement supposée. Lors du développement du nouveau test, deux critères en particulier sont mis en avant : le comportement face à des échantillons de petite taille et la robustesse. Malgré le potentiel de chacun des tests proposés dans la littérature pour la détection des ruptures dans les séries multivariées, une limitation majeure et commune pour tous ces tests est qu'ils soient basés sur l'hypothèse que le type de la copule reste invariant. Le nouveau test proposé permet de relaxer cette hypothèse. En effet, il permet de détecter la rupture aussi bien dans le paramètre que dans le type de la copule. L'intérêt du nouveau test réside également dans sa formulation générale et flexible qui permet d'analyser l'homogénéité de différents types de série hydrologique. D'autre part, le nouveau test est capable de détecter les changements dans les séries faiblement ou fortement dépendantes et de distinguer entre plusieurs types de copules, même pour les échantillons de petite taille. La section suivante décrit brièvement les aspects théoriques du test proposé.

4.1.2. Méthodologie

Le but du nouveau test proposé est de vérifier si une série des données est homogène ou non. Formellement, il permet de tester l'hypothèse nulle $H_{0,C}: \exists C$ tel que la dépendance entre $X_1 \dots X_n$ est décrite par C contre l'hypothèse alternative $H_{1,C}: \exists$ un indice k et $\left\{ \begin{array}{l} \exists C_1 \text{ tel que la dépendance entre } X_1 \dots X_k \text{ est décrite par } C_1 \\ \exists C_2 \text{ tel que la dépendance entre } X_{k+1} \dots X_n \text{ est décrite par } C_2 \end{array} \right.$

La statistique du test proposé est basée sur les L-moments multivariés. Afin qu'elle soit indépendante de l'unité de mesure, la statistique proposée est définie à partir des *rapports* de L-moments par

$$T_n = \max_{1 \leq k \leq n} \|\Lambda_k - \Lambda_{n-k}\| \quad (1.7)$$

$$\text{tel que } \Lambda = \begin{bmatrix} \tau_2^C & \tau_3^C & \tau_4^C \end{bmatrix}, \tau_i^C = \begin{bmatrix} \tau_{i[11]}^C & \tau_{i[21]}^C \\ \tau_{i[12]}^C & \tau_{i[22]}^C \end{bmatrix} \quad i = 2, 3, 4 \quad \tau_{k[12]}^C = \frac{\lambda_{k[12]}^C}{\lambda_{2[2]}^C}, \text{ pour } k \geq 3$$

et $\tau_{2[12]}^C = \frac{\lambda_{2[12]}^C}{\lambda_{1[2]}^C}$ et $\lambda_{k[12]}^C$: L-moment d'ordre k de la copule C .

En général, une copule peut être caractérisée par ses quatre premiers L-moments (e.g. Brahimi *et al.*, 2015). Par conséquent, les deux copules C_1 et C_2 sont différentes si la différence entre les quatre premiers L-moments dépasse un seuil limite. L'utilisation des L-moments est avantageuse puisque ces derniers permettent de caractériser plusieurs copules. De plus, leurs estimateurs sont moins biaisés, comparés aux autres mesures de dépendance telles que le rho de Spearman ou le tau de Kendall (e.g. Brahimi *et al.*, 2015; Brahimi and Necir, 2012; Capéraà and Genest, 1993). Comme la distribution de la statistique du test n'est pas explicite et afin de contrôler l'erreur de première espèce à un seuil de signification α donné, la méthode d'auto échantillonnage « *bootstrap* » est employée (e.g. Genest and Rémillard, 2008; Gombay and Horvath, 1999).

Afin de montrer la bonne performance de la statistique proposée, une étude de simulation est réalisée. Le but de cette étude de simulation est d'évaluer la puissance du test proposé dans un contexte hydrologique. La performance peut être affectée par la taille de l'échantillon traité, l'amplitude du changement et la distribution multivariée de la série étudiée. Par conséquent, une étude de sensibilité à ces facteurs est effectuée. Par la suite, pour comparaison, les tests existants qui semblaient les plus prometteurs ont été sélectionnés sur la base de leur potentiel d'applicabilité aux séries hydrologiques. En outre, deux cas d'étude sont présentés pour démontrer l'aspect pratique du nouveau test selon des contraintes hydrologiques. Bien que le test proposé soit valable pour une diversité de variables hydrologiques, dans le cadre de la présente thèse, nous nous concentrons sur l'étude des caractéristiques des crues, particulièrement le débit Q et le volume V .

4.2. Nouvelle méthode d'estimation des paramètres de copules mixtes

Cette section décrit l'objectif et la méthodologie développés dans l'article [A2] intitulé « *Meta-heuristic estimation method for mixture copula models* ». Les détails méthodologiques ainsi que les résultats sont décrits dans le chapitre 3 correspondant à l'article [A2].

4.2.1. Description de l'objectif

Comme discuté dans la section 3.1, il est important de vérifier l'homogénéité d'une série avant d'entamer l'étape de modélisation (objectif de l'article [A1]). Si l'hétérogénéité de la série est statistiquement significative, l'utilisation systématique d'un modèle de copule simple peut avoir des conséquences négatives sur l'estimation des risques associés à un événement donné. Par conséquent, il est pertinent de considérer un modèle plus adéquat qui tient compte de cette hétérogénéité. Ainsi, nous proposons un modèle mixte constitué d'un mélange de copules (avec éventuellement des mélanges au niveau des marges). Une fois les copules formant les composantes sont connues, la principale difficulté réside dans l'estimation des paramètres du modèle. D'où le deuxième objectif spécifique, qui consiste à proposer une nouvelle méthode d'estimation des paramètres du modèle de copules mixtes. En plus de combler les lacunes liées aux estimateurs existants, ce nouvel estimateur a l'avantage d'être basé sur une optimisation simple évitant de nombreux problèmes d'instabilité tout en conservant de bonnes qualités de convergence. Une brève description de l'approche méthodologique est présentée dans la section suivante.

4.2.2. Méthodologie

En statistique en général et en *AF* en particulier, la méthode du maximum de vraisemblance (Maximum Likelihood, *ML*) est la plus utilisée pour l'estimation des paramètres en raison de son optimalité et de sa flexibilité (e.g. Caudill and Acharya, 1998; El Adlouni *et al.*, 2007; Leytham,

1984). De même, cette méthode a été adaptée pour l'estimation des paramètres de la copule (e.g. Genest *et al.*, 1995; Genest and Rivest, 1993; Salvadori and De Michele, 2010). Or, dans le cas de la copule, la vraisemblance est fonction des marges, qui sont généralement inconnues. Par conséquent, afin de contourner le problème d'interférence des marges dans l'estimation des paramètres de la copule, la méthode *MPL* a été proposée. Pour résoudre de façon efficiente le problème de maximisation associé à la méthode *MPL*, nous proposons les algorithmes génétiques (*AG*) comme solution d'optimisation. En effet, contrairement à l'estimateur obtenu par la méthode *EM*, en plus de contourner les problématiques soulevées dans la section 3.2, l'estimateur *AG* est consistant, dans le sens où il converge toujours vers le maximum global. De plus, l'originalité des *AG* est de considérer non pas une valeur initiale des paramètres, mais une population initiale des paramètres, qu'il s'agit de faire évoluer jusqu'à la convergence de tous les paramètres la composant vers le maximum global.

Des simulations sont réalisées pour tester la performance des *AG* en termes d'estimation, mais également pour les comparer avec la méthode *EM*, dans un contexte hydrologique. À partir de séries synthétiques, trois statistiques mesurant la performance de chaque estimateur ont été considérées : l'erreur relative, le biais relatif et l'erreur quadratique relative. Les simulations ont été réalisées à partir de plusieurs modèles et scénarios afin d'obtenir un ensemble complet décrivant ce qui pourrait se produire lors d'une *AF* des variables hydrologiques.

4.3. Nouveaux tests d'adéquation pour les copules multiparamètres

Dans cette section, nous présentons l'objectif de l'article [A3] « *Multivariate L-moment based tests for copula selection, with hydrometeorological applications* ». Une brève description de la méthodologie est également présentée. Les détails méthodologiques ainsi que les résultats sont décrits dans le chapitre 4 correspondant à l'article.

4.3.1. Description de l'objectif

En hydrologie, les copules multiparamètres revêtent certains intérêts. En effet, elles offrent plus de flexibilité et permettent de caractériser plus d'un type de dépendance simultanément. De plus, leur structure flexible a l'avantage de décrire des structures de dépendance très diverses, notamment quand les coefficients de queue inférieure et de queue supérieure sont différents. Par ailleurs, si une hétérogénéité est détectée dans la série de données (Objectif de l'article [A1]), tel que discuté dans la section 4.2, il est important d'incorporer cette hétérogénéité en utilisant un modèle de mélange des copules (Objectif de l'article [A2]). Le résultat de ce modèle de mélange des copules peut être considéré comme une copule multiparamètre. Or, la structure de dépendance réelle des données observées est inconnue. Lors de la modélisation de cette dernière, différents candidats de modèles se présentent. Éventuellement, le plus grand souci lors du choix d'un modèle est sa capacité de bien s'ajuster aux données. À ce titre, les tests d'adéquation sont des outils indispensables.

Bien que dans la littérature, il existe plusieurs tests d'adéquation, comme discuté dans la section 3.3, ces derniers souffrent de plusieurs limites. À cet égard, le troisième objectif spécifique de la thèse vise à développer des tests d'adéquation spécifiques pour les copules multiparamètres dont le but consiste à sélectionner la copule la plus adéquate à modéliser la dépendance dans un phénomène donné. En plus de contourner les limites des tests existants, les nouveaux tests ont l'avantage d'être tout à fait adaptés aux caractéristiques des variables hydrologiques. Le paragraphe suivant explique brièvement l'approche méthodologique de cet objectif.

4.3.2. Méthodologie

Les tests d'adéquation proposés permettent de vérifier si une copule multiparamètre s'ajuste bien aux séries de données. Formellement, ils permettent de tester l'hypothèse nulle $H_0: C \in$

$\{C_{\vec{\theta}}, \vec{\theta} \in \mathbb{R}^p\}$ contre l'alternative $H_1: C \notin \{C_{\vec{\theta}}, \vec{\theta} \in \mathbb{R}^p\}$. Pour les mêmes raisons que celles évoquées dans la section 4.1.2, les nouveaux tests sont basés sur les L-moments multivariés. La première statistique du test consiste à comparer une distance entre les rapports des L-moments multivariés de la copule choisie et ceux empiriques. Si la distance entre les deux rapports de L-moments est grande, alors le modèle est rejeté. La deuxième statistique est une extension de celle développée par Berg and Quessy (2009). Cette dernière est basée sur la comparaison de deux estimateurs des paramètres de la copule. Dans le test proposé ici, les deux estimateurs choisis sont ceux obtenus par les méthodes des L-moments et du maximum de pseudo-vraisemblance. Ce choix est motivé par les propriétés de consistance et de convergence de ces deux estimateurs.

La performance des tests proposés est évaluée dans un premier temps sur des séries simulées. Une analyse de sensibilité par rapport aux facteurs pouvant affecter la puissance des tests est effectuée. Ensuite, la puissance des tests d'adéquation proposés est comparée à celle de quelques tests existants et récents. Enfin, les tests développés dans cette thèse sont appliqués aux séries hydrologiques réelles.

Deuxième partie: les articles scientifiques

Chapitre 2. Homogeneity testing of multivariate hydrological records, using multivariate copula L-moments

Titre en français : Détection de l'hétérogénéité dans les séries hydrologiques multivariées

Auteurs :

Imene Ben Nasr Étudiante au doctorat à l'Institut National de la Recherche Scientifique (Centre Eau Terre Environnement), Université du Québec, 490 Rue de la Couronne, Québec, QC, Canada, G1K 9A9, Téléphone: (418) 261-6864, courriel: imene.ben_nasr@ete.inrs.ca

Fateh Chebana Professeur, Institut national de la recherche scientifique (Centre Eau Terre Environnement), Université du Québec, 490 Rue de la Couronne, Québec, QC, Canada, G1K 9A9, Téléphone: (418) 654-2542, courriel: Fateh.Chebana@ete.inrs.ca

A1. Ben Nasr, I. and F. Chebana (2019). "Homogeneity testing of multivariate hydrological records, using multivariate copula L-moments." *Advances in Water Resources* 134:103449. DOI: <https://doi.org/10.1016/j.advwatres.2019.103449>.

Dans cet article, I. Ben Nasr a proposé un nouveau test statistique pour la détection de l'hétérogénéité dans les séries hydrologiques multivariées, basées sur les L-moments. Le développement et l'écriture ont été effectués par I. Ben Nasr. Le co-auteur F. Chebana a commenté et révisé la version finale du manuscrit.

Abstract

Recently, there have been an increasing number of studies dealing with change detection in multivariate series. However, a major drawback with most of the currently used methods is the lack of flexibility. Indeed, these methods are only able to detect changes in the strength of dependence assuming invariant shape of the dependence structure. However, under a changing climate, the shape of dependence might change as well. Furthermore, in the multivariate setting, heterogeneities can occur in the margins and/or in the dependence structure. Most of the existing approaches for multivariate change detection deal with the whole distribution. In this paper, we propose a novel statistical test for multivariate heterogeneity detection, based on copula and multivariate L-moments. A simulation study is conducted to evaluate the performance of the proposed test and to compare it with those of existing tests. Results indicate that the proposed test has an interesting power especially when the dependence strength remains invariant, with power ranging between 45 and 93% whereas for the existing tests the power is lower than 14% in this case. An application to a real data set is also provided. Results show the ability of the proposed test to discriminate homogeneous and inhomogeneous series.

Keywords: Multivariate, Homogeneity testing, Stationarity, Multivariate L-moments, Copula, Flood.

1. Introduction

Extreme events are of high importance in design of hydraulic structures, water resources management and flood insurance. Consequently, their estimation should be as reliable as possible. For this purpose, hydrological frequency analysis (*HFA*) is the most used statistical tool to model behaviour of hydro-meteorological variables (e.g. Hamed and Rao, 1999). These events are characterized by a number of correlated variables (e.g. Chebana, 2013; Salvadori and De Michele, 2004). For example, floods can be described by the peak, volume and duration (e.g. Requena *et al.*, 2013b; Shiau *et al.*, 2006). Therefore, a multivariate framework that takes into account the dependence between different variables is of fundamental importance (e.g. Requena *et al.*, 2013b; Zhang and Singh, 2012). Copulas have been proposed as a key tool to model the dependencies between hydrological variables in the context of multivariate HFA (e.g. Chebana, 2013; Genest and Chebana, 2016; Salvadori and De Michele, 2004; Salvadori *et al.*, 2007).

HFA is composed of four main steps, as summarized in table 2.1 (e.g. Chebana *et al.*, 2010). Checking basic assumptions (homogeneity, stationarity and serial independence) is an important step since it has a significant impact on the other subsequent steps. Ignoring this step could lead to catastrophic consequences in case of an underestimation of these events and waste of valuable economic resources in case of their overestimation (e.g. Bender *et al.*, 2014). In hydrological multivariate setting, checking for trend-freeness is recently treated by Chebana *et al.* (2013), while homogeneity testing has attracted less attention despite its importance. Moreover, departures from this assumption may provide misleading results. Furthermore, if unaccounted for, these inhomogeneities could have a large impact on the outcome of the estimation and fitting of the probability distribution (e.g. Ouarda *et al.*, 2014; Seidou and Ouarda, 2007).

Table 2.1. HFA steps in the univariate and multivariate frameworks

HFA steps	Univariate	Multivariate
i) Exploratory analysis and outlier detection	Large body of literature	Very sparse literature
ii) Checking assumptions		a) Very sparse literature: the aim of the present paper
a) Homogeneity	Large body of literature (e.g. Alila and	b) Very sparse literature (Chebana <i>et al.</i> , 2013)
b) Stationarity	Mtiraoui, 2002; Yue <i>et al.</i> , 2002)	c) Sparse literature (e.g. Fan <i>et al.</i> , 2017; Taskinen <i>et al.</i> , 2005)
c) Independence		
iii) Modeling and estimation	Large body of literature (e.g. Khaliq <i>et al.</i> , 2006)	Large literature (e.g. Chebana and Ouarda, 2011; Salvadori and De Michele, 2010; Vittal <i>et al.</i> , 2015)
iv) Risk assessment and analysis	Large body of literature (e.g. Rao and Hamed, 2000)	Growing literature (e.g. Fan <i>et al.</i> , 2015; Sarhadi <i>et al.</i> , 2016)

Data are considered as homogenous if they have the same underlying distribution, with invariant statistical characteristics (e.g. Gilroy and McCuen, 2012; Lund *et al.*, 2007; Reeves *et al.*, 2007). Hence, it means that data do not exhibit a mix of several samples from different populations. It is necessary to note that this definition does not involve time. Consequently, as pointed out by Hamed and Rao (1999), the chronological order is not important for a sample to be homogenous. The lack of homogeneity may be caused by natural or anthropogenic actions on physical environment, such as deforestation, dam construction and urbanization among others (e.g. Burn and Elnur, 2002). Furthermore, annual flood series might be produced by more than one hydrometeorological process such as flood-producing storms, rainfall and snowmelt floods (e.g. Ouarda and El-Adlouni, 2011; Smith *et al.*, 2011). Moreover, the regime of a river downstream from the confluence of two sub-watersheds with very different hydrological behaviors is a good example of the lack of homogeneity.

In the literature, the most common way for testing homogeneity, in both univariate and multivariate context, is to detect abrupt change-point (e.g. Das and Umamahesh, 2017; He *et al.*, 2016; Quessy, 2019; Ribes *et al.*, 2017; Sadegh *et al.*, 2015). In fact, homogeneity and change-point are two

related problems but the latter is more challenging. As illustrated by Figure 2.1, homogeneity problems have two or more pre-specified sub-groups of data to be tested (without necessary a sequential aspect of the data). However, in the case of a change-point problem, the appropriate groupings of the data are unknown and data has to be segmented (in a sequential manner) into sub-groups that will be tested for homogeneity. Homogeneity problem can be considered as a change-point testing problem by using preliminary step, such as classifying or clustering data (e.g. Lung-Yut-Fong *et al.*, 2015). Hence, in this paper, homogeneity problem is treated as a change-point detection (*CPD*) problem. To deal with the *CPD* problem, formal parametric and nonparametric testing approaches have been proposed in the literature. In this study, only nonparametric approaches are considered since they do not require prior choice of marginal and joint distributions which makes them more appropriate in *HFA* (e.g. Vittal *et al.*, 2015). For general comparison between parametric and nonparametric frameworks, the reader is referred for instance to Santhosh and Srinivas (2013) and Oja and Randles (2004).

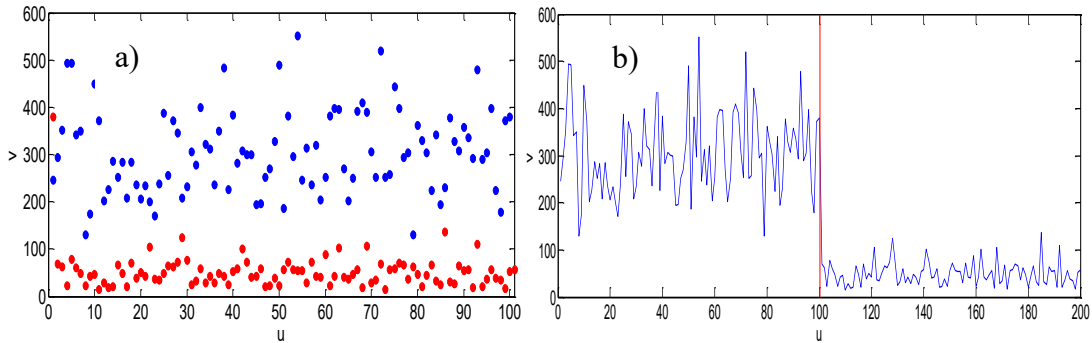


Figure 2.1. Difference between a) homogeneity and b) change-point testing problem

The hydro-meteorological literature abounds with studies dealing with homogeneity testing of univariate distributions (e.g. Beaulieu *et al.*, 2009; Ehsanzadeh *et al.*, 2011; Reeves *et al.*, 2007; Seidou and Ouarda, 2007). However, the multivariate setting has been scarcely addressed (e.g. Chebana *et al.*, 2017). Homogeneity of multivariate hydrological series can affect marginal

distributions as well as the dependence structure (e.g. Guerfi *et al.*, 2015; Holmes *et al.*, 2013; Xiong *et al.*, 2015). Therefore, it is of high importance to distinguish between these two types of homogeneities in order to gain deep physical insights behind the causes of the eventual changes in the data (e.g. Mazouz *et al.*, 2012).

A general class of tests, based on multivariate statistical depth, has been introduced and evaluated for the hydrological context in Chebana *et al.* (2017). These tests allow detecting an overall shift in the series. However, they do not allow distinguishing if the shift is in margins or in the dependence structure. To deal with this problem, a few multivariate tests have been recently proposed to detect inhomogeneities in the dependence structure (e.g. Holmes *et al.*, 2013; Quessy *et al.*, 2013; Xiong *et al.*, 2015). Nonetheless, as summarized in Table 2.2, these tests entail some drawbacks. For instance, some tests are designed to detect changes only in the parameter of the copula and not able to detect changes in the shape of dependence such as the test proposed by Bouzebda and Keziou (2013). This is an important limitation since under lack of homogeneity, the copula itself could change. Moreover, some of these tests have optimal performance under normality. This is a major drawback since hydrological series are generally non-normally distributed. To avoid these limitations, another class of tests based on Kendall and empirical copula process are proposed (e.g. Kojadinovic *et al.*, 2016; Quessy *et al.*, 2013). Unfortunately, these tests are not suitable for small sized samples. Indeed, these tests allow detecting change on long data length (for $n > 100$), which is not usually available in practice. In addition, a change in the copula may occur even if Kendall process remains invariant such as the case of extreme-value copulas (e.g. Genest and Favre, 2007; Vogel and Fried, 2015). Moreover, changes captured by these tests usually tend to appear in the middle of series, whereas heterogeneities may also occur at the beginning or at the end of the records.

Table 2.2. Overview of existing tests for change detection in the dependence structure

Test based on	Reference	Designed to detect	Limitations
Kendall's τ	Dehling <i>et al.</i> (2017); Quessy <i>et al.</i> (2013)	Change in the dependence strength	• Unable to detect change in the dependence when τ remains constant. • Powerful when change occurs on the middle of the series. • Power increases as the variance of Kendall's τ decreases. • Lack power against alternatives for moderate and small sample sizes.
Spearman's ρ	Kojadinovic <i>et al.</i> (2016); Wied <i>et al.</i> (2014)	Change in the dependence strength	• Choice of a Kernel smoothing parameter and a bandwidth. • Have no power against alternatives when ρ does not change or only very little. • Powerful for large sample size.
Empirical process	Bücher and Volgushev (2013); Holmes <i>et al.</i> (2013)	Change in the dependence structure	• Not sensitive to a change in the copula when margins are homogenous. • Do not differentiate the type of change: marginal or in the dependence. • Powerful when the type of copula remains constant.
Likelihood-ratios	Bouzebda and Keziou (2013)	Change in the dependence strength	• Optimal under normality. • Detect only change on the dependence strength: copula parameter. • Invariant copula type • Parametric copula specification. • Time consuming

In order to avoid the aforementioned drawbacks, the objective of the present paper is to develop a new test for homogeneity of the dependence structure. The developed test is based on multivariate L-moments given their attractive desirable properties for hydrological applications. Indeed, they provide a summary and a description of the properties and shapes of a multivariate distribution. This makes them particularly useful in parameter estimation and hypothesis testing. Furthermore, since multivariate (also univariate) L-moments are much less biased than classical moments, they are used as meaningful replacements of classical moments in a wide variety of applications, mainly in hydrology, climatology and meteorology analysis (e.g. Chebana and Ouarda, 2007; Kysely and Picek, 2007). A performance evaluation of the proposed test, based on Monte-Carlo simulations,

is presented under hydrological constraints. A case study is also performed to illustrate an application of the developed test on hydrological data.

Aside from the introduction, the remainder of the paper is organized as follows. A brief review of definitions and some notions, related to the developed test, is presented in section 2. Section 3 presents the development of the proposed test for homogeneity testing. The simulation study to evaluate the performance of the test is presented in Section 4. Section 5 illustrates an application of the developed test on hydrological data. The conclusions of the study and several perspectives are reported in section 6.

2. Theoretical background

In this section, we tackle the statistical tools used in the current paper. A brief description of the statistical background of homogeneity testing problem is provided herein.

2.1. Multivariate homogeneity testing problem

As mentioned earlier, homogeneity is a fundamental hypothesis in *HFA* since its absence may affect the statistical inference to be undertaken. Homogeneity of a sample postulates that all data are generated by the same physical process (e.g. Hamed and Rao, 1999; Peterson *et al.*, 1998). This implies, therefore, checking that all data have an invariant probability distribution (e.g. Ehsanzadeh *et al.*, 2011; Ribes *et al.*, 2017).

Shift detection on location and/or dispersion parameters is the commonly used form of testing the homogeneity of the hydrological series (e.g. Hamed and Rao, 1999; Kundzewicz and Robson, 2004; Li and Liu, 2004; Seidou *et al.*, 2007). As pointed out by Rougé *et al.* (2013) and Nayak and Villarini (2016), this formulation of the homogeneity problem is based on the sequential time of the event. However, heterogeneity can be present without importance of the sequential aspect. As

a consequence, throughout this paper, we focus on the homogeneity in form of distributional change. In statistical terms, this means that not only some characteristics of the distribution change abruptly, but also the distribution itself changes (e.g. Milly *et al.*, 2015; Serinaldi *et al.*, 2018).

In the multivariate *HFA*, as pointed out by Vezzoli *et al.* (2017), heterogeneities may occur in marginal distributions, dependence structure or both. Over the past few decades, a large number of techniques have been developed to detect heterogeneities in marginal distribution. However, as far as we know, the homogeneity of the dependence structure (copula type) has been scarcely addressed (e.g. Guerfi *et al.*, 2015; Xiong *et al.*, 2015). Since hydrological time series increasingly exhibit non-stationarity due to natural and anthropogenic changes, this issue is worthy of attention in hydrologic designs.

To define the homogeneity testing problem, the null hypothesis H_0 is that, when arbitrarily splitting the sample in two subsamples, there is no change in the distribution function of each subsample. Statistically, let $(X_i)_{i=1,\dots,n}$ be a d -variate random variables of size n , with marginal cumulative distribution function (*m.c.d.f.s*) $(F_j)_{j=1\dots d}$ and multivariate joint cumulative distribution function G . We denote $x_i = (x_i^1, \dots, x_i^d)'$ the observation from X_i . Under H_0 , it is assumed that all observations are homogeneous and thus have the same distribution G . The alternative hypothesis H_1 assumes G to exhibit inhomogeneity *i.e.* G is time-varying and is no longer constant.

In agreement with Sklar (1959) 's theorem, there exists a copula C such that, for all x_i , we have:

$$G(x_1, \dots, x_d) = C(F_1(x_1), \dots, F_d(x_d)) \quad (2.1)$$

Regarding the links between copulas and multivariate joint distributions, the null hypothesis H_0 can be expressed formally as $H_0 = H_{0,m} \cap H_{0,c}$ against the alternative $H_1 = H_{1,m} \cup H_{1,c}$.

$H_{0,m}$ and $H_{0,C}$ (respectively $H_{1,m}$ and $H_{1,C}$) are the null (respectively the alternative) hypothesis about the homogeneity of marginal distributions and copula. Homogeneity testing of marginal distributions and copula are expressed respectively as:

$$H_{0,m}: \exists F_1, \dots, F_d \text{ such that } X_1 \dots X_n \text{ have m.c.d.f.s } F_1, \dots, F_d$$

$$H_{1,m}: \exists \text{ at least an index } s \text{ and } \begin{cases} \exists F_{1,1}, \dots, F_{1,d} \text{ such that } X_1 \dots X_s \text{ have m.c.d.f.s } F_{1,1}, \dots, F_{1,d} \\ \exists F_{2,1}, \dots, F_{2,d} \text{ such that } X_{s+1} \dots X_n \text{ have m.c.d.f.s } F_{2,1}, \dots, F_{2,d} \end{cases}$$

$$H_{0,C}: \exists C \text{ such that } X_1 \dots X_n \text{ have copula } C$$

$$H_{1,C}: \exists \text{ at least an index } s \text{ and } \begin{cases} \exists C_1 \text{ such that } X_1 \dots X_s \text{ have copula } C_1 \\ \exists C_2 \text{ such that } X_{s+1} \dots X_n \text{ have copula } C_2 \end{cases}$$

As mentioned above, this paper focuses on the homogeneity of the copula structure. Therefore, the proposed statistic attempts to test $H_{0,C}$ against $H_{1,C}$.

In the following, for the sake of convenience and brevity, only the bivariate case is considered.

2.2. Multivariate copula L-moments

The multivariate L-moment was introduced by Serfling and Xiao (2007) in both parametric and nonparametric schemes. Multivariate L-moments provide a summary and a description of the properties and shapes of a distribution. This makes them particularly useful in parameter estimation and hypothesis testing. Since multivariate (also univariate) L-moments are much less biased than classical moments, they are used as meaningful replacements of classical moments in a wide variety of applications, especially for hydrological frequency analysis (e.g. Ben Nasr and Chebana, 2019b; Chebana and Ouarda, 2007; Hosking and Wallis, 1993).

The multivariate L-moments are defined as bellow (Serfling and Xiao, 2007).

$$\lambda_{k[ij]} = cov(X^{(i)}, P_{k-1}^*(F_j(X_j))) \quad (2.2)$$

where COV is the covariance, k is the order of the multivariate L-moment, P_{k-1}^* is the shifted Legendre polynomial (Chang and Wang, 1983), X^i is a random variable with marginal distribution F_i for $i = 1, 2$. Brahimi *et al.* (2015) introduced the copula's multivariate L-moments as

$$\lambda_{k[12]}^C = \int_0^1 \int_0^1 (C(u_1, u_2) - u_1 u_2) du_1 dP_k^*(u_2) \quad (2.3)$$

The sample version of the k^{th} copula multivariate L-moments is defined in term of pseudo-observations (u_i, v_i) as (Brahimi *et al.*, 2015)

$$\lambda_{k[12]}^C = \frac{1}{n} \sum_{i=1}^n u_i P_{k-1}^*(v_i) \quad (2.4)$$

where pseudo-observations $u_i = \frac{R_i}{n+1}$; $v_i = \frac{S_i}{n+1}$ and R_i is the rank of X_i^1 among $X_1^1, \dots, X_n^1$; S_i is the rank of the concomitant $X_{i:n}^{[12]}$ among $X_i^2, \dots, X_n^2$. To obtain $X^{[12]}$, we first sort X^2 in the ascending order as $X_{1:n}^2 < X_{2:n}^2 \dots < X_{n:n}^2$. Then, the concomitant $X_{i:n}^{[12]}$ corresponds to X_i^1 paired with $X_{i:n}^2$ (Serfling and Xiao, 2007). The copula multivariate L-moments coefficient *ratios* are defined as

$$\tau_{k[12]}^C = \frac{\lambda_{k[12]}^C}{\lambda_{2[12]}^C}, \text{ for } k \geq 3 \text{ and } \tau_{2[12]}^C = \frac{\lambda_{2[12]}^C}{\lambda_{1[2]}^C} \quad (2.5)$$

3. Multivariate L-moments-based copula homogeneity test

Recall that the null hypothesis H_0 of no change in the dependence structure of the d -variate variables $(X_i)_{i=1..n}$ postulates that the copula C is invariant against the alternative H_1 that there is k such that $X_1, \dots, X_k$ have copula C_1 and $X_{k+1}, \dots, X_n$ have copula C_2 . If H_0 is rejected, k is the so called change-point in dependence structure. The proposed test statistic is then given by

$$T_n = \max_{1 < k < n} \|\Lambda_k - \Lambda_{n-k}\| \quad (2.6)$$

where Λ_k and Λ_{n-k} are matrices whose elements are multivariate L-moments ratios of order 2, 3 and 4 of the series before and after the point k , respectively; $\|\bullet\|$ is the Euclidean norm.

The dimension of the matrices Λ_k and Λ_{n-k} is (2×6) . Each matrix of multivariate L-moments

coefficients Λ is written as $\Lambda = \begin{bmatrix} \tau_2^C & \tau_3^C & \tau_4^C \end{bmatrix}$ where $\tau_i^C = \begin{bmatrix} \tau_{i[11]}^C & \tau_{i[21]}^C \\ \tau_{i[12]}^C & \tau_{i[22]}^C \end{bmatrix}$ $i = 2, 3, 4$.

The null hypothesis H_0 should be rejected for a large value of the test statistic T_n . The change-point is, therefore, estimated by

$$k^* = \underset{1 \leq k \leq n}{\operatorname{argmax}} \|\Lambda_k - \Lambda_{n-k}\| \quad (2.7)$$

The asymptotic distribution of the proposed statistic is out of the scope of the present paper and could be the object of future work. In addition, asymptotic results could not be appropriate for hydrological application given the short sample sizes usually encountered. Hence, the bootstrap procedure is used to estimate *p-values* of the proposed test (Good, 2004; Holmes *et al.*, 2013).

The idea of the proposed test is inspired by the tests proposed by Quessy *et al.* (2013) and Kojadinovic *et al.* (2016). However, it allows avoiding some drawbacks of the traditional tests. In fact, three key features distinguish our test from the existing ones (given in Table 2.2). First, the tests by Quessy *et al.* (2013) and Kojadinovic *et al.* (2016) are designed to detect only the change in the dependence strength (Kendall or Spearman coefficients) and hence, assuming the same copula shape before and after the change-point. In contrast, the proposed test allows for detecting change on the shape of copula even if the Kendall's τ or Spearman's ρ remains invariant. Second, these two tests are not well adapted to particularities of hydrological series, where they have low power for small sample sizes. The proposed test is based on multivariate L-moments, which are accurate for small sample size usually encountered in *HFA*. Third, unlike the likelihood ratios test (Bouzebda and Keziou, 2013), the proposed test is nonparametric which is robust against

distribution miss-specification and respects the fact that the test should be performed prior the modeling step in *HFA*. Semi-parametric tests should perform well under the correct model specifications, but inference based on miss-specified models is not well studied.

The heterogeneity of a data series may be due to differences in any feature of the distribution. In particular, in the hydrological literature, the L-covariation matrix Λ_2 , represents a measure of the dispersion of the distribution (e.g. Chebana and Ouarda, 2007). Besides, the L-coskewness Λ_3 is an important descriptor of the copula structure. In particular, Λ_3 measures the difference between the upper and lower tails, and hence measures the asymmetric shape of the copula (e.g. Müller *et al.*, 2017). Likewise, Λ_4 gives an idea about the kurtosis of multivariate distribution by measuring the difference between the typical spread in the tails and the typical spread in the center of the multivariate distribution. Consequently, taken together, the L-covariation, L-coskewness and L-cokurtosis matrices allow detecting change on the dependence structure, from a low-L-moments perspective. Note also that the first order multivariate L-comoment Λ_1 is not included in the proposed test statistic since it represents the componentwise mean vector and hence, does not affect the dependence structure (Serfling and Xiao, 2007). Note also that the proposed statistic T_n is a rank-based since the matrices of the L-moments ratios are expressed as function of the pseudo-observations, which are the best summary of the joint behavior of the random pairs and hence not influenced by the margins (e.g. Genest and Chebana, 2017; Kao and Govindaraju, 2010; Kao and Govindaraju, 2008; Salvadori and De Michele, 2004). In addition, the proposed statistic has several advantages including:

- a. Simple formulas are available in terms of the ranks of the observations, hence it can be applied without assuming any prior distribution about data;

- b. The proposed statistic has the advantage that it makes no assumption on the type of copula being assessed. Hence, it can be applied for any copula type;

The decision rule regarding the acceptance or rejection of the null hypothesis is based on the *p-value*. To evaluate the associated *p-value*, resampling methods, such as permutation and bootstrapping, are used (e.g. Good, 2004).

In the present work since the focus is on adequately detect and estimate the location where the change occurs, we focus on the statistic T_n , which combine all L-moments ratios together. Indeed, as explained by Serfling and Xiao (2007), distinct distributions generate distinct series of L-moments. Furthermore, Kjeldsen and Prosdocimi (2015a) argued that L-moments of different orders capture sharply different population features and provide an improvement in the ability to detect the underlying distribution when compared to the performance of the one-dimensional L-moment. Hence, in the same vein, using a combination of different L-moments ratios is more convenient for detecting change in the dependence and permits a high probability of discrimination between different copulas.

4. Simulation study

The purpose of the simulation study is to evaluate the performance of the proposed multivariate homogeneity test. To this end, we consider practical cases commonly encountered in hydrological applications. Despite the validity of the proposed test for different hydrological events such as floods, rain storms and droughts, in the present paper, we restrict attention to the performance of the proposed test when dealing with floods.

4.1. Simulation design

In practice, flood events are described by several correlated variables, namely the peak Q , the volume V and duration D (e.g. Aissia *et al.*, 2012; Fu and Butler, 2014; Shiau *et al.*, 2006). Since much of the available literature on flood events deals with dependence between Q and V , these two variables are considered in this study. Generated data are used to test the ability of the proposed test to detect heterogeneities, as well as its ability to identify the position when the change occurs.

Data are generated through simulations from representative and most used copulas in hydrometeorology analyses (e.g. Salvadori and De Michele, 2010; Salvadori *et al.*, 2007; Zhang and Singh, 2006). Therefore, three different classes of copula family are used to model the dependence structure between Q and V , namely, Archimedean, Extreme-Value and Elliptical. To generate data from a given copula, we applied the procedure proposed by Nelsen (2006), based on the conditional distribution method. In particular, the employed copulas are Frank and Clayton (Archimedean), Gumbel and Galambos (Extreme-Value) and Gaussian (elliptical). It is worth noting that the aforementioned copulas are considered under different scenarios to generate heterogeneous data. Therefore, data are mixture of two copulas. To avoid confusion, note that dependence shape stands for copula type as described by a specific dependence structure.

Hereafter, different scenarios are considered.

- a) Homogenous data: all data are generated from the same copula;
- b) Heterogeneous on the dependence shape: data are generated from two different copulas with same dependence strength;
- c) Heterogeneous on the dependence strength: data are generated from the same copula with different dependence strengths;

- d) Heterogeneous on the dependence structure: data are generated from two different copulas with different dependence strengths;

According to previous studies, Q and V can be marginally fitted by the Gumbel distribution (e.g. Chebana and Ouarda, 2007; Requena *et al.*, 2013b; Villarini and Smith, 2010). Consequently, we consider the Gumbel distribution, with scale and location parameters denoted respectively by σ and β , as marginal for both Q and V . The corresponding parameters of the considered distribution are those obtained from the data series of the Skootamatta basin in Ontario, Canada, used for previous simulation studies by Chebana *et al.* (2017) and Chebana and Ouarda (2009). In this study, we consider a change in the location of Q and V simultaneously. Hence, for a change located at the middle of the series, the corresponding parameters of the Gumbel distribution before the change are defined as: $\sigma_Q = 15.85$, $\beta_Q = 51.85$, $\sigma_V = 300.22$, $\beta_V = 1239.8$. These parameters are also used to generate homogenous series of Q and V . Then, for generating data after the change, we consider amplitude of change $(\delta_Q, \delta_V) = (40\%, 40\%)$. For more details the reader can refer to Chebana *et al.* (2017).

The performance of a homogeneity test (either univariate or multivariate) could be affected by various factors, specifically the record length, the dependence strength, the change-point position and the copula type (e.g. Bouzebda and Keziou, 2013; Holmes *et al.*, 2013; Quessy *et al.*, 2013). Hence, a sensitivity analysis of the performance of the proposed test is performed, regarding different factors.

First, the sample size n is a relevant factor to the homogeneity of a series. Hydrological series are generally characterised by small sized samples. Hence, the assessment of the behaviour of the proposed test was performed under several sample sizes $n = 30, 50, 100$. The values of n are

selected on the basis of situations commonly encountered in flood frequency analysis (see series in Barth *et al.* (2017) and Santhosh and Srinivas (2013)).

The performance of the homogeneity test can also be affected by the dependence strength measured by Kendall's τ . The test statistic should ideally be able to distinguish between heterogeneous subsamples having the same τ . Thus, the smaller the influence of τ on the statistic the better the test will be. This analysis is performed by varying the dependence strength τ . Since for most flood events the dependence strength is typically between 0.3 and 0.8 (e.g. Requena *et al.*, 2013a; Zhang and Singh, 2007b), three values of $\tau = 0.2, 0.6, 0.8$, corresponding to weak, moderate, and strong dependence, respectively, are considered. The parameter of the copula is defined to match the corresponding range of dependence, as described through Kendall's τ (e.g. Nelsen, 2006; Vandenberghe *et al.*, 2010). The estimation of the parameter of the copula is achieved by using Kendall's τ inversion method described in Nelsen (2006) and Joe (2014).

Besides sample size and dependence strength, the location of the change is an important factor that may influence the performance of the homogeneity test in detecting heterogeneity in the dependence structure. In the literature (e.g. Dehling *et al.*, 2017; Nayak and Villarini, 2016; Ribes *et al.*, 2017; Xiong *et al.*, 2015), findings suggest that, in general, homogeneity test is more powerful when the change occurs far from the beginning or the end of the series. Furthermore, some univariate tests are preferable if the change occurs in the middle of the series. In order to determine the ability of the proposed test to accurately detect departure from homogeneity according to when a change occurs, three locations are considered. Hence, a change is taken to occur at location $s = \frac{n}{4}, \frac{n}{2}, \frac{3n}{4}$. A diagram of the simulation study is shown in Figure 2.2.

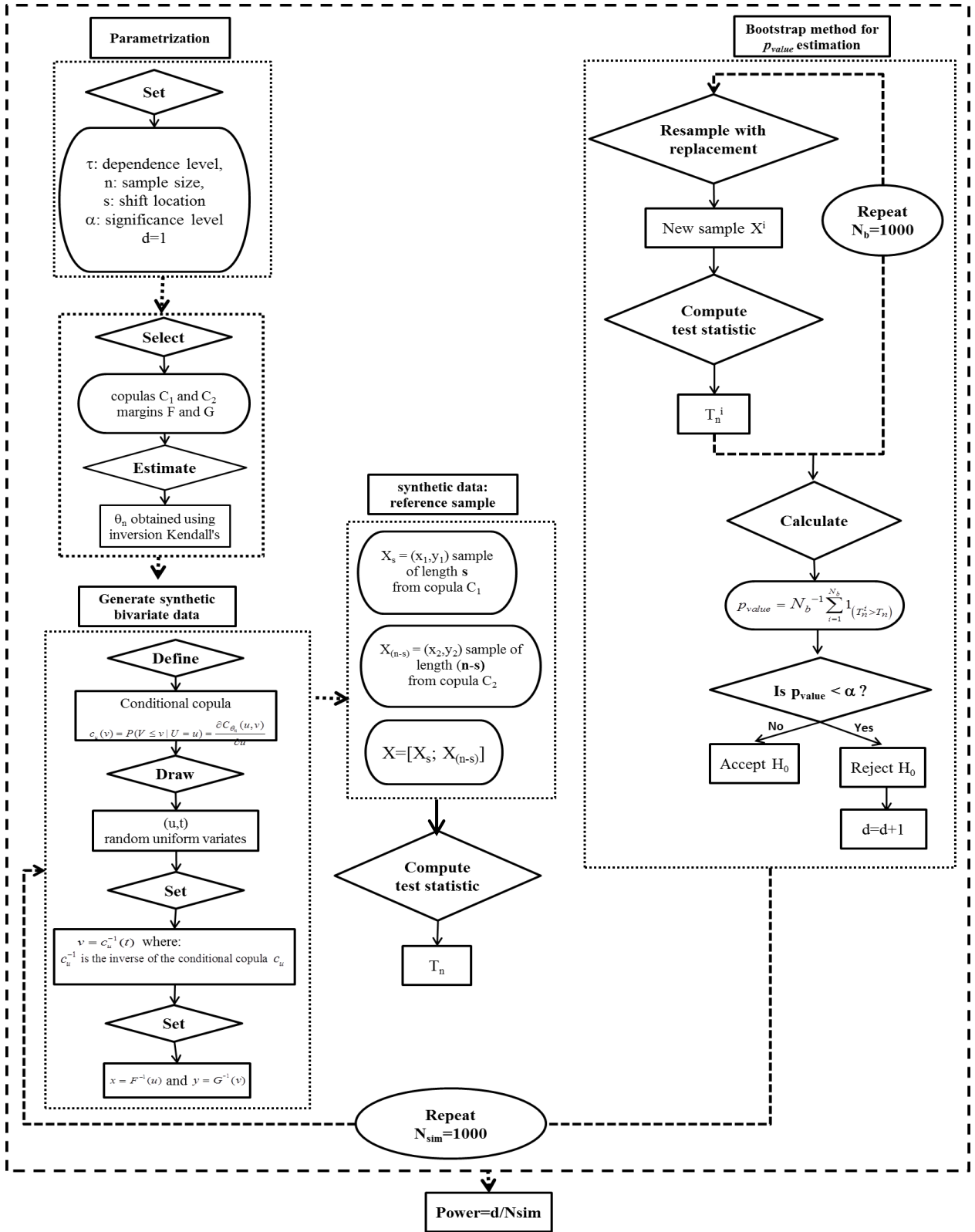


Figure 2.2. Diagram of the simulation study

The flowchart is made up of three different parts. The different rectangles indicate the different building blocks of our methodology: synthetic data generation, p_{value} estimation and power computation. For all scenarios, $N_{sim}=1000$ samples are generated which is typically used in homogeneity testing simulations (e.g. Holmes *et al.*, 2013; Xiong *et al.*, 2015). For each generated sample, the proposed statistic T_n is computed on a moving time window of size n_w , chosen arbitrary by the user (e.g. Bender *et al.*, 2014). The time window length n_w must be chosen in a way to be neither too large nor too small in order to provide stationary variables within the windows and a reliable estimation of multivariate L-moments. Then, the time window is shifted by one observation each time step and the corresponding first four multivariate L-moments ratios are calculated. In the present study, it is chosen as $n_w=25, 15$ and 10 respectively for large, moderate and small samples.

As is well-known in the hypothesis testing, two features are of interest, namely, the nominal level (α) and the power ($1-\beta$) (Good, 2004). The first is that the null hypothesis is rejected while it is true. The second is that the null hypothesis is accepted when the alternative is true. These two features can be estimated, using simulations study, by computing the rejection rates. In the present study, we fix $\alpha=5\%$, which is frequently used in the hydrological literature regarding hypothesis testing (e.g. Xie *et al.*, 2014; Xiong *et al.*, 2015).

4.2. Simulation results

To evaluate the performance of the proposed test, it is important to analyze a large number of simulated series representing a variety of situations. In this section, results obtained by the application of the proposed test are presented.

4.2.1. Nominal level evaluation

The first desirable property of a test is to hold the nominal level close to the significance level, under the null hypothesis. Table 2.3 reports the estimates $\hat{\alpha}$ of α for the proposed test, for different sample sizes and dependence strengths. First, it can be seen from Table 2.3 that the proposed statistic T_n , when considering the different copula types investigated in this study, is not sensitive to the copula type regarding the estimation of α for a given dependence strength and sample size. Indeed, $\hat{\alpha}$ is between 4.7 and 5.2%, for $\tau=0.8$ and $n=100$, for different copula type. It is also interesting to note, from Table 2.3, that the estimates $\hat{\alpha}$ improve when the sample size increases, e.g. for the Clayton copula and $\tau=0.8$, it decreases from 8.2% when $n=30$ to 4.9% when $n=100$. Indeed, this is likely due to the fact that the uncertainty of estimates of the multivariate L-moments decreases with the sample size (e.g. Brahimi *et al.*, 2015). It is worth noting that Holmes *et al.* (2013); Quessy *et al.* (2013); and Xiong *et al.* (2015) also found similar effect of the sample size on the nominal level of their tests, which are designed to detect the change in the strength of dependence (copula parameters).

Table 2.3. First type error estimates (%) by the proposed test for different sample sizes and different dependence strengths

Copula under H_0		$\tau=0.2$			$\tau=0.6$			$\tau=0.8$		
		n=30	n=50	n=100	n=30	n=50	n=100	n=30	n=50	n=100
Archimedean	Clayton	11.2	7.5	5.2	11.2	5.9	5.0	8.2	5.8	4.9
	Frank	10.8	7.2	6.3	10.5	6.2	5.2	8.1	5.4	4.7
Extreme-Value	Gumbel	12.1	6.1	5.5	10.8	5.9	5.1	7.2	5.6	5.0
	Galambos	11.3	7.4	6.2	9.9	6.9	5.8	7.6	6.1	5.2
Elliptical	Gaussian	11.7	8.3	6.0	12.3	6.4	6.0	9.6	5.9	5.1

This table presents the rejection rates of the null hypothesis by the proposed test at significance level $\alpha=5\%$, for different scenarios generated under the null hypothesis.

The effect of the dependence strength is also shown in Table 2.3. It is worth noting that the test is not highly sensitive to the dependence strength when sample size is large enough. Indeed, for different $\tau=0.2, 0.6$ and 0.8 , $\hat{\alpha}$ is close to 5%, when the sample size $n=100$. It shows that in general

the estimates $\hat{\alpha}$ decrease with the dependence strength, to reach the significance level $\alpha=5\%$. As an example, for the Clayton copula with $n=50$, it decreases from 7.5% to 5.8% when τ increases from 0.2 to 0.8. Likewise, by applying the test of Bouzebda and Keziou (2013) (see Table 2.2) in an hydrological case, Xiong *et al.* (2015) found that the highest dependence strength is, the better estimated $\hat{\alpha}$ will be. Hence, copula with a strong dependence is more easily to be detected.

4.2.2. Power evaluation

In this section, the power of the proposed test in detecting the change in the dependence structure is studied. Results, for different sample sizes and dependence strengths, are displayed in Table 2.4.

Table 2.4. Power estimates (%) of the proposed statistic T_n

CP*	$(\tau_b, \tau_a)^{**}$	n=30			n=50			n=100		
		(Cl,Fr)	(G,Fr)	(Gm,Fr)	(Cl,Fr)	(G,Fr)	(Gm,Fr)	(Cl,Fr)	(G,Fr)	(Gm,Fr)
n/4	(0.4,0.4)	48.2	45.5	49.2	76.0	74.1	77.7	86.4	88.2	90.4
	(0.2, 0.6)	61.0	60.9	63.7	86.8	87.4	88.2	90.6	90.2	91.0
	(0.4, 0.8)	55.5	56.1	59.0	85.8	84.8	86.2	89.8	90.0	90.2
n/2	(0.4,0.4)	63.0	63.2	63.4	85.8	86.8	87.4	90.8	90.6	91.8
	(0.2, 0.6)	67.0	69.2	69.4	90.6	89.6	90.0	93.6	92.6	94.0
	(0.4, 0.8)	65.8	66.4	67.6	89.8	89.0	89.0	90.8	91.2	92.4
3n/4	(0.4,0.4)	44.7	44.0	47.0	77.6	74.8	79.9	89.6	88.4	88.8
	(0.2, 0.6)	64.0	63.8	64.4	86.6	85.6	84.8	91.0	89.0	93.0
	(0.4, 0.8)	52.4	58.4	62.4	76.8	80.9	84.0	90.6	88.2	91.4
n/4	(0.4,0.4)	43.4	48.4	48.8	73.2	75.7	71.6	82.2	85.7	83.4
	(0.2, 0.6)	61.6	64.0	65.0	88.6	89.6	89.0	90.8	90.2	91.6
	(0.4, 0.8)	52.4	54.9	56.8	75.2	77.2	78.4	85.9	89.2	89.6
n/2	(0.4,0.4)	58.8	55.6	59.4	86.8	89.4	90.0	91.2	91.4	92.0
	(0.2, 0.6)	66.2	67.4	68.8	91.0	90.6	92.2	92.0	93.9	94.2
	(0.4, 0.8)	63.8	66.4	66.2	90.0	89.0	91.4	91.6	92.8	94.0
3n/4	(0.4,0.4)	46.2	47.4	48.9	74.5	75.3	73.2	83.5	81.1	82.3
	(0.2, 0.6)	64.4	61.8	66.2	88.4	87.8	87.4	91.2	90.6	93.4
	(0.4, 0.8)	61.8	60.0	64.0	86.6	85.8	86.4	90.8	89.9	90.4

CP is the location of the given change-point; (τ_b, τ_a) corresponds to the dependence strength before and after the change; Cl, F, G, Gm stand for Clayton, Frank, Gaussian and Gumbel copula, respectively. This table presents the power of the proposed test at significance level $\alpha=5\%$, for different scenarios. Grey color stands for homogenous margins.

From Table 2.4, we can see that the power of the test is sensitive to the position where a change occurs ($s=n/2$ or $s=n/4$). Indeed, for a given sample size and dependence strength, the test has best powers when the change occurs on the middle of the series. For example, when data are a mixture of Frank and Gumbel copulas with $n=50$ and dependence strength $(0.4, 0.4)$, the power increases from 77% when $s=n/4$ to 87.4 % when $s= n/2$. However, when the sample size increases ($n=100$), the test is less sensitive to the change-point position. In addition, overall, copula type seems to have little influence on the powers of the proposed test. Therefore, no significant differences were found between powers when considering different copula types. For example, for $n=100$ and dependence strength $\tau= (0.2, 0.6)$ with a change located at $s=n/2$, powers are between 92.6% and 93.6% when data are mixture of Clayton-Frank copulas and Gaussian-Frank copulas, respectively.

Through Table 2.4, we can see a variation in the powers regarding the dependence strength τ . In fact, the proposed test performs clearly better when the amplitude of the jump in Kendall's τ increases. As an example, for a series of length $n=50$, generated from a mixture of Gaussian and Frank copula and a change located at $s=n/4$, the test power increases from 74.1% for $\tau= (0.4, 0.4)$ to 87.4% when $\tau= (0.2, 0.6)$. Hence, in this case, the change associated to larger change range in dependence strength, is more easily to be detected. Similar findings are reported in the literature regarding change detection on dependence strength (e.g. Xiong *et al.*, 2015). Results from Table 2.4 show that the proposed statistic is sensitive to the amplitude of the change in the dependence strength. However, it is less sensitive to the dependence strength. In fact, the power of the proposed test for $\tau= (0.2, 0.6)$ is larger than for $\tau= (0.4, 0.8)$. Thus, the larger is the variance of τ , the better is the power. Moreover, it is important to emphasise that the proposed statistic T_n performs well in detecting change in the copula shape. Namely, the test is able to detect the heterogeneity when the copula structure before and after the change point are different and the Kendall's τ remains the

same. In fact, for the same dependence strength, the test has a high power in departure from H_0 . For instance, for $n=100$ and a mixture of Frank and Gumbel copula with $\tau=(0.4, 0.4)$, the power estimates are between 82.3 and 91.8%.

Table 2.4 also provides insights about the effect of sample size on the power of the proposed test. There is a clear trend of increasing power as sample size increases. In fact, this tendency can be a consequence of the fact that when sample size is small ($n=30$), the rejection of the null hypothesis is easier since the dependence shape is more difficult to distinguish. Moreover, this can be explained by the higher uncertainty in estimating the L-moments matrices for small samples. In fact, for sample size $n=30$, the size of subsample to estimate the L-moments matrices is equal to 10, which is too short to have a reliable estimator, especially in the multivariate case (Brahimi *et al.*, 2015). These results agree with the findings of other studies, in which the power increases with the sample size (e.g. Quessy *et al.*, 2013; Xiong *et al.*, 2015).

4.2.3. Comparison

As mentioned in the literature review, existing multivariate homogeneity tests deal only with the change on the dependence strength. A comparison of the performance between the proposed test and the classical ones is made through simulation studies as summarized in Table 2.5 and 2.6.

Table 2.5. Power comparison between different tests when considering the same copula (here Gumbel) before and after the change

CP^*	$(\tau_b, \tau_a)^{**}$	$n=30$			$n=50$			$n=100$		
		T_n	K_n	S_n	T_n	K_n	S_n	T_n	K_n	S_n
$n/4$	(0.2, 0.6)	36.5	20.1	13.2	47.8	33.0	32.6	56.9	53.9	40.7
	(0.4, 0.8)	32.5	19.4	12.8	39.4	32.0	31.0	55.7	52.1	41.6
$n/2$	(0.2, 0.6)	48.0	24.1	19.2	59.1	51.4	35.8	67.0	61.6	47.2
	(0.4, 0.8)	40.7	22.9	15.2	58.0	53.7	34.2	68.9	59.0	49.0
$3n/4$	(0.2, 0.6)	34.7	17.6	11.3	41.2	37.2	30.1	64.7	53.6	44.6
	(0.4, 0.8)	30.3	12.8	10.0	38.5	31.9	30.4	61.1	48.0	41.8

CP is the location of the given change-point; (τ_b, τ_a) corresponds to the dependence strength before and after the change. This table presents the power of different tests at significance level $\alpha=5\%$, for different scenarios, considering homogenous margins. T_n : the proposed statistic, K_n : Kendall's τ test and S_n : the likelihood-ratio test.

Table 2.5 provides corresponding results when the copula structure remains invariant before and after the change point (here Gumbel copula is considered), for different sample sizes, dependence strengths and change-point positions. According to Table 2.5, the power estimates of all statistics increase when n becomes larger. This agrees with similar findings by Dehling *et al.* (2017); Lung-Yut-Fong *et al.* (2015); and Xiong *et al.* (2015). Since the proposed test is based on multivariate L-moments, it is appropriate for hydrological applications, where the record lengths are typically short. As one can see from Table 2.5, power estimates of the proposed test, range respectively, between 30.3 and 48% for $n=30$ and between 38.5 and 59.1% for $n=50$. This can be explained by the fact that multivariate L-moments are robust and suffer less from the effects of sampling variability (Brahimi *et al.*, 2015). Nevertheless, for the same design, the two classical statistics are not able to detect the change point for small sized-samples. The associated powers of these tests are between 10.1 and 24.1% for $n=30$, and between 30.1 and 53.7% for $n=50$. These results agree with the findings of other studies in which Xiong *et al.* (2015) suggested that the copula-likelihood ratio test performs poorly for a relatively small sample size ($n = 50$). In addition, change-point in small samples proved to be harder to capture by the test proposed by Quessy *et al.* (2013). The performance of the proposed test in small samples shows the relevance of the introduced approach.

One of the advantages of the proposed multivariate homogeneity test might be its ability to detect not only change in the dependence strength but also in copula type. These interesting features of the proposed test will now be compared to the existing tests. Results regarding the application of the proposed test as well as classical tests, to detect heterogeneities when the dependence structure (copula type) before and after the change are different, are summarized in Table 2.6. It is notable from Table 2.6 that the likelihood ratio test is not able to detect the change in the dependence structure. In all cases, the power estimates are less than 15%. This is unsurprising since this test is

conceived to detect only the change in the dependence strength. On the other hand, the Kendall's τ test has good performance for higher change in the amplitude of dependence strength. Indeed, for sample size $n=100$ and a change point located at the middle of the series, power estimates of the Kendall's tau test are 53.7% when $\tau_{before}=0.2$ and $\tau_{after}=0.6$ and 6.7% when $\tau_{before}=\tau_{after}=0.4$.

Table 2.6. Power comparison between different tests, when considering different copulas before and after the change (here Gumbel and Frank)

CP*	$(\tau_b, \tau_a)^{**}$	n=30			n=50			n=100		
		T _n	K _n	S _n	T _n	K _n	S _n	T _n	K _n	S _n
n/4	(0.4,0.4)	45.5	7.6	6.2	74.1	7.1	6.3	88.2	4.9	12.5
	(0.2, 0.6)	60.9	18.2	6.5	87.4	35.8	7.2	90.2	45.9	11.9
	(0.4, 0.8)	56.1	19.5	4.9	84.8	32.6	7.5	90.0	40.4	10.8
n/2	(0.4,0.4)	63.2	7.2	8.0	86.8	7.0	8.6	90.6	6.7	13.9
	(0.2, 0.6)	69.2	32.6	6.7	89.6	47.6	7.5	92.6	53.7	10.8
	(0.4, 0.8)	66.4	28.0	5.2	89.0	43.3	6.0	91.2	52.0	12.7
3n/4	(0.4,0.4)	44.0	6.5	5.0	74.8	5.0	5.2	88.4	6.1	12.1
	(0.2, 0.6)	63.8	26.8	7.0	85.6	41.2	6.7	89.0	43.5	15.7
	(0.4, 0.8)	58.4	21.7	7.9	80.9	39.4	6.4	88.2	41.4	10.1

CP is the location of the given change-point; (τ_b, τ_a) corresponds to the dependence strength before and after the change. This table presents the power of different tests at significance level $\alpha=5\%$, for different scenarios, considering homogenous margins. T_n: the proposed statistic, K_n: Kendall's τ test and S_n: the likelihood-ratio test.

This means that high power estimates are obtained for largest difference in the dependence strength before and after the change-point. Furthermore, it is apparent from Table 2.6 that the proposed test can capture change in copula type, when the dependence strength remains invariant which is not the case of the aforementioned tests. Powers of the proposed test range between 88 and 91% when keeping the same Kendall's τ before and after the change-point, for sample size $n=100$. However, for the other tests, powers are between 4.9 and 13.9%.

It is also important to note that the proposed test has approximately the same computation time as classical tests. Table 2.7 shows the computation time for performing each test 100 times, for the same scenario. It is shown that there is a huge gain in computation time with the proposed L-moments test and the Kendall's τ test, compared to the Likelihood ratios.

Table 2.7. Execution time for 100 evaluations

Sample size	Proposed L-moments based test	Kendall's τ based test	Likelihood ratio test
$n=30$	2'25''	2'52''	5'11''
$n=50$	4'04''	4'49''	8'27''
$n=100$	5'46''	6'29''	12'44''

Considered copulas before and after the change point are Gumbel and Frank, with the same Kendall's $\tau=0.4$

Besides the power and the first type error estimation, an additional important feature of a change point-detection test is the ability to well-positioning the change. In the univariate framework, a change-point is considered well-positioned when it is located within ± 2 years of the true position (e.g. Beaulieu *et al.*, 2009). We use the same threshold for the multivariate case. So far this paper has focused on change detection in the copula structure, the following section will compare the detection skills of the proposed test statistic (T_n) to the Kendall-based test (K_n) proposed by Quessy *et al.* (2013). Here, the likelihood ratio test is not considered since it is based in the assumption of invariant copula structure. Figure 2.3 compares the percentage of the well-identified change-point, obtained by each test, among the detected change point, for each scenario combining different sample sizes, dependence ranges and copula types. It is apparent from this figure that the Kendall test is not able to well positioned the detected change point for small sample size ($n=30, 50$). For the proposed statistic T_n , the position of the detected change is well-identified for different dependence range. However, we can see that the performance is an increasing function of the sample size.

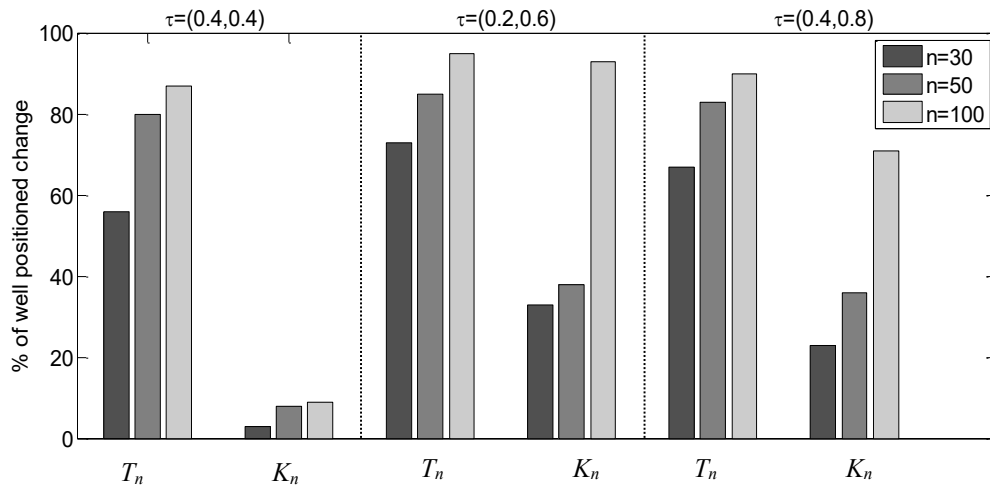


Figure 2.3. Percentage of well positioned change detected by different statistics for different scenarios

5. Applications to Real Hydrological Data

In this section, we apply the test on two real-world hydrological data. The purpose of these applications is to assess the appropriateness of the proposed test for practical use.

The first data series correspond to the Romaine station with natural flow regimes, located in the Cote Nord Region of the province of Quebec, Canada. It has been considered previously for change-point studies by Chebana *et al.* (2017). The second data series corresponds to the Ankang hydrological station at the Upper Hanjiang River, China. The same series is used on the paper of Xiong *et al.* (2015), who demonstrated the existence of a change on the copula parameter. General information about the used data are given in Table 2.8. For each station, flood characteristics: $Peak(Q)$ and $Volume(V)$, are extracted from daily streamflow series (Aissia *et al.*, 2012).

Table 2.8. General information about the Romaine and Ankang stations

Station Name	Latitude (°)	Longitude (°)	Period of record (years)	Catchment Area (km ²)
Romaine	50.18	-63.37	1957-2013	12922
Ankang	32.71	109.03	1975-2011	38600

For the Romaine station, the p -value, associated to the proposed multivariate statistic is 2%. This result suggests the existence of a change on the dependence structure of (Q, V) . Furthermore, the

estimated change-point is located at 1984, which corresponds to the detected change on the univariate variables Q and V (Chebana et al., 2017). In the same way, for the Ankang station, the multivariate change-point is located at 1986, with a corresponding p -value equal to 4%. Note that this year corresponds, approximately, to the beginning of water storage at the Ankang reservoir, located about 30 km upstream the Ankang hydrological station (Jiang et al., 2015).

In order to validate results regarding detected change-point, different copulas are fitted to the entire series and to each subsample before and after the detected change, similarly to the set up used by Salvadori et al. (2018). The goodness-of-fit (GOF) test selected to formally identify which copulas represent properly the dependence structure of the data is that based on the empirical copula and the Cramér-von Mises statistic (e.g. Genest et al., 2009). Then, to select the best fitting copula, AIC criterion is used (e.g. Laio et al., 2009). Results are summarised in Table 2.9.

Table 2.9. Results of the goodness-of-fit test and model selection criteria, for copula selection

Station	Romaine						Ankang					
	All data		Before 1984		After 1984		All data		Before 1986		After 1986	
	P_{value}	AIC	P_{value}	AIC	P_{value}	AIC	P_{value}	AIC	P_{value}	AIC	P_{value}	AIC
Clayton	0.003	-19.83	0.006	-3.30	0.007	-7.84	0.794	-163.26	0.030	-63.50	0.105	-73.44
Frank	0.039	-24.28	0.027	-8.13	0.165	-9.55	0.661	-151.03	0.402	-77.64	0.113	-71.29
Gumbel	0.146	-25.21	0.131	-10.25	0.034	-6.58	$2 \cdot 10^{-4}$	-138.34	0.208	-71.47	0.025	-67.63
Normal	0.447	-28.63	0.045	-8.94	0.794	-9.91	0.138	-157.07	0.140	-79.70	0.062	-69.90
Hüssler-Reiss	0.427	-27.30	0.436	-10.98	0.065	-7.84	$7 \cdot 10^{-4}$	-134.57	0.049	-72.05	0.020	-63.08
Galambos	0.077	-26.24	0.212	-10.61	0.080	-7.21	10^{-4}	-137.92	0.030	-71.58	$3 \cdot 10^{-4}$	-67.26

This table presents the p -value of the GOF test and AIC criterion for different copula model fitted to the time series without and with change-point, using the maximum pseudo-likelihood method. Gray color corresponds to the accepted copula by the GOF test at the significance level 5%. Bold character corresponds to the best fitting copula.

A set of univariate distributions are also considered as possible candidates for fitting the studied variables: the Weibull, Gumbel, Log-normal, and Gamma distributions. The Anderson-Darling GOF test (Sinclair et al., 1990) is used to discard those distributions that cannot represent the observed univariate data. Results are presented in Table 2.10.

Table 2.10. Results of the goodness-of-fit test and model selection criteria, for univariate distribution selection

Station		Romaine		Ankang	
Variable	C.P.	Distribution	pvalue	AIC	pvalue AIC
Q	all data	Weibull	0.012	900	0.192 154
		Gumbel	0.030	918	0.086 159
		Gamma	0.175 825	0.135 157	
		Lognormal	0.003	859	0.030 164
	Before C.P.	Weibull	0.154 401	0.208 64	
		Gumbel	0.054	403	0.114 66
		Gamma	0.037	420	0.140 65
		Lognormal	0.003	418	0.096 68
	After C.P.	Weibull	0.041	434	0.180 93
		Gumbel	0.011	432	0.020 97
		Gamma	0.091 426	0.127 94	
		Lognormal	0.002	432	0.044 97
V	all data	Weibull	0.030	908	0.169 1219
		Gumbel	0.002	910	0.148 1223
		Gamma	0.179	903	0.034 1250
		Lognormal	0.183 900	0.051 1233	
	Before C.P.	Weibull	0.041	437	0.155 509
		Gumbel	0.031	437	0.140 511
		Gamma	0.058	435	0.017 531
		Lognormal	0.079 434	0.158 512	
	After C.P.	Weibull	0.122 474	0.020 737	
		Gumbel	0.104	479	0.068 734
		Gamma	0.112	476	0.003 741
		Lognormal	0.120	475	0.280 729

The best fitting univariate distribution and copula structure are presented in Table 2.11. For the Romaine station, under the homogeneity hypothesis, the normal copula ($\theta = 0.685$) is considered as the best copula that represents the relation between Q and V . However, under the heterogeneity hypothesis, the copula that best fits the observed data before the change-point 1984 is the Hüssler-Reiss copula whereas after 1984, it is the normal copula. It is worthwhile noting that the two selected copulas, before and after 1984, have very close dependence strength ($\tau_{before}=0.42$, $\tau_{after}=0.41$). This result illustrates the capability of the proposed test to distinguish between two copulas with same dependence strength, as shown in simulation study.

Table 2.11. Selected univariate distribution and copula

Station	C.P.	Variable	Distribution	Parameters
Romaine	all data	Q	<i>Gamma</i>	$(\beta=12.27, \eta=124.81)$
		V	<i>Normal</i>	$(\mu=3846.81, \sigma=839.23)$
		<i>Copula</i>	<i>Normal</i>	$\theta=0.68, \tau=0.45$
	Before 1984	Q	<i>Weibull</i>	$(\beta=5.17, \eta=1873.37)$
		V	<i>Normal</i>	$(\mu=4236.93, \sigma=710.01)$
		<i>Copula</i>	<i>Hüssler-Reiss</i>	$\theta=1.60, \tau=0.42$
	After 1984	Q	<i>Gamma</i>	$(\beta=12.11, \eta=112.86)$
		V	<i>Weibull</i>	$(\beta=6.21, \text{scale}=4536.38)$
		<i>Copula</i>	<i>Normal</i>	$\theta=0.64, \tau=0.42$
Ankang	all data	Q	<i>Weibull</i>	$(\beta=1.99, \eta=9.95 \cdot 10^3)$
		V	<i>Weibull</i>	$(\beta=2.02, \text{scale}=1.86 \cdot 10^9)$
		<i>Copula</i>	<i>Clayton</i>	$\theta=8.64, \tau=0.85$
	Before 1986	Q	<i>Weibull</i>	$(\beta=2.66, \eta=1.12 \cdot 10^4)$
		V	<i>Weibull</i>	$(\beta=2.71, \eta=2.14 \cdot 10^9)$
		<i>Copula</i>	<i>Normal</i>	$\theta=0.91, \tau=0.81$
	After 1986	Q	<i>Weibull</i>	$(\beta=1.51, \eta=7.51 \cdot 10^3)$
		V	<i>Lognormal</i>	$(\log \mu=20.70, \log \sigma=0.72)$
		<i>Copula</i>	<i>Clayton</i>	$\theta=12.19, \tau=0.91$

This table presents the best fitting univariate distribution as well as the copula structure. β stands for the scale parameter and η the shape parameter of the distribution. θ is the copula parameter and τ is the dependence strength.

For the Ankang station, the Clayton copula, with a dependence parameter $\theta=8.64$ (obtained by the inversion Kendall's tau method), is selected as the best fitting copula under homogeneity. However, when considering the change-point, before 1986, normal copula is chosen as the best candidate reaching dependence strength 0.81 and after 1986, the dependence increases to 0.91 with a change on the shape of the copula, from normal to Clayton.

As pointed out through this study, it is of high importance to detect changes and to take into account these changes in the modeling steps. Indeed, if there is any heterogeneity in the data, the statistical parameters of the model, estimated under homogeneity, do not reflect the physics of the phenomenon, and interpretations made are erroneous.

Since different combinations of (u, v) pairs can lead to the same $C(u, v)$ value, it is common to express the joint probability $C(u, v)$ via contours of equal probability (i.e., copula probability level

curves) where $C(u, v) = p$ with $0 < p < 1$. Associated copula probability level curves, before and after the detected change-point, are presented in Figure 2.4.

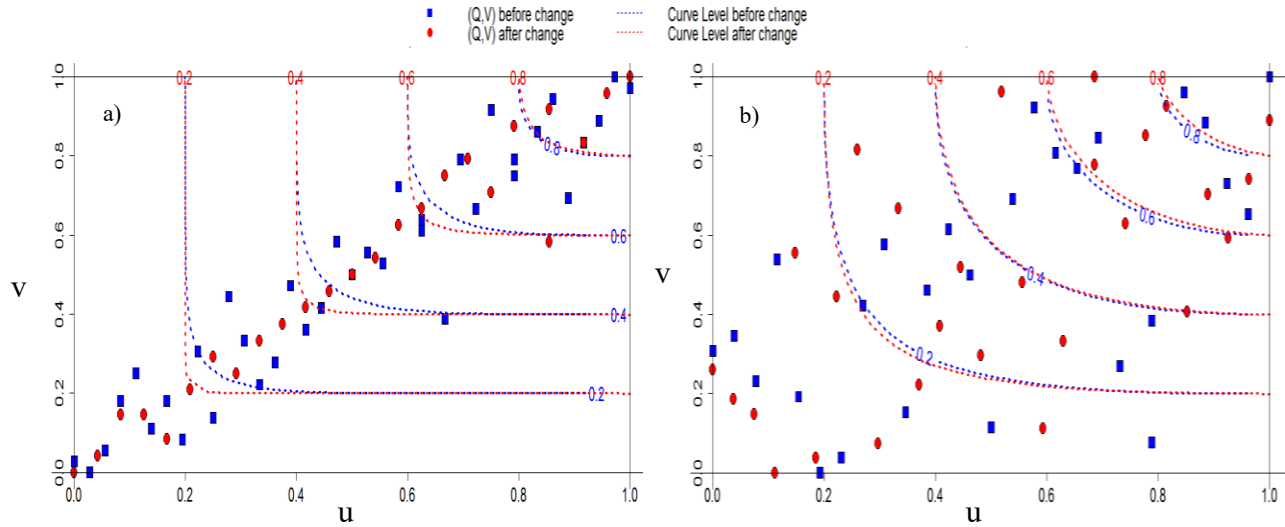


Figure 2.4. Comparison between level curves of the fitted copula before and after the change for a) the Ankang and b) Romaine stations

As shown in Figure 2.4a, in the case of the Ankang station, the level curves corresponding to different given joint probabilities change dramatically before and after the change in the copula structure. However, for the Romaine station, the change is less visible and the difference between level curves before and after the change-point is very small. These results are likely to be related to the fact that the Kendall's τ , for the Romaine station, is almost the same before and after the change ($\tau_{before}=0.42$, $\tau_{after}=0.41$). Only for illustration purposes, Figure 2.5 shows that the difference between level curves, associated to normal and hüssler-Reiss copulas (the same selected copulas for this station), are more pronounced when there is a change on the Kendall's τ (dependence strength). Note also that the change of the level curves is more accentuated when the change of the copula is coupled with the change in marginal distributions (Figure 2.6). As a consequence, this might affect the results of the multivariate frequency analysis especially in terms of underestimation of the associated risk.

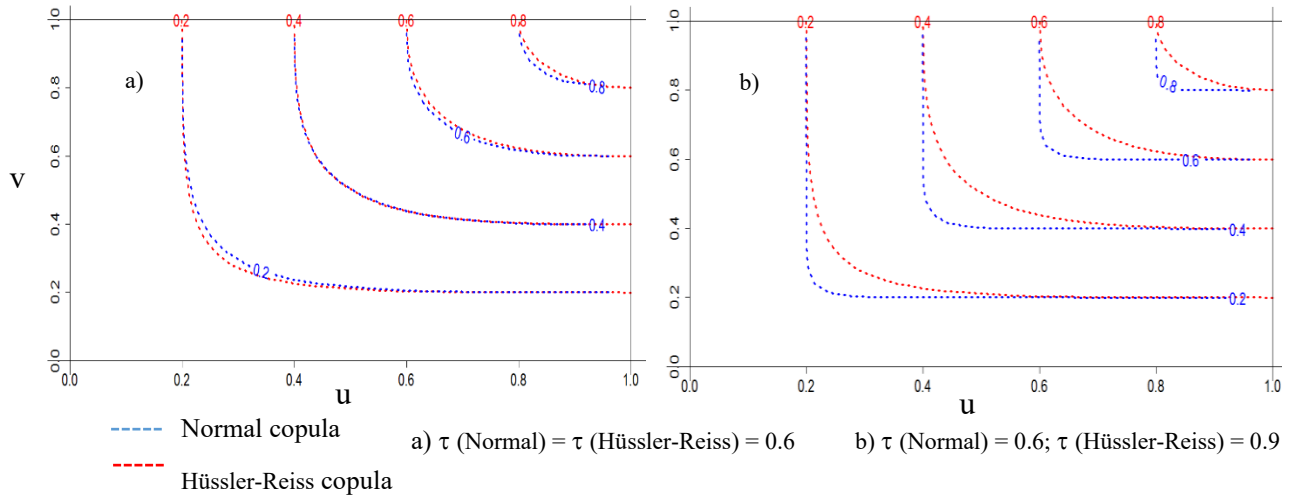


Figure 2.5. Level curves of the normal and hüslerReiss copulas for a) same Kendall's τ and b) different Kendall's τ

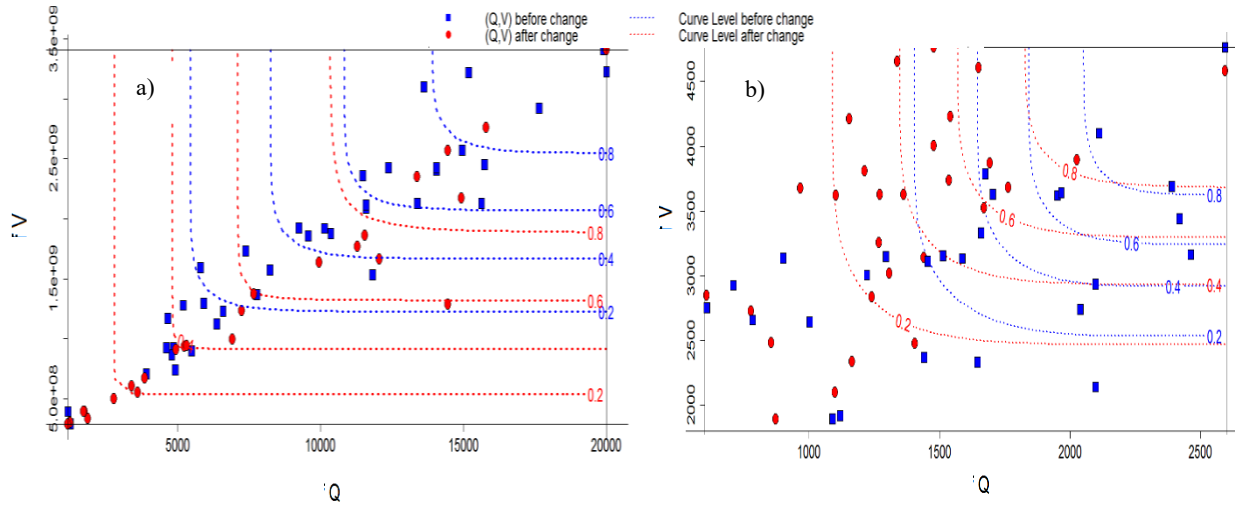


Figure 2.6. Comparison between level curves of the multivariate distribution before and after the change, for a) the Ankang and b) Romaine stations

6. Conclusion

While there are a large number of homogeneity tests in the literature, only few tests deal with multivariate series. Furthermore, little attention has been paid to the homogeneity of the whole dependence structure. Existing tests detect the change-point in the strength of the dependence through the parameter of a given copula, assuming the same type of copula.

In this paper, a new test for homogeneity of multivariate dependence structure is proposed, and it is adapted to the specifications of the hydrological context. The proposed test is based on the

multivariate L-moments. In order to evaluate its performance, simulations are considered as well as real-world flood case studies. Results show very interesting performances in term of first type error and power estimations. More precisely, findings reveal that this test is suitable for small-sized hydrological series and for series with same dependence strength. The application of the proposed test to real data shows its ability to detect and well-position the change-points, especially cases where the dependence strength remains the same whereas the change concerns the whole copula structure. The proposed test can be considered for a wide range of hydrological, water resources and climatological applications.

Chapitre 3. Meta-heuristic estimation method for mixture copula models

Titre en français : Méthode méta-heuristique pour l'estimation des paramètres de copules mixtes

Auteurs :

Imene Ben Nasr Étudiante au doctorat à l'Institut National de la Recherche Scientifique (Centre Eau Terre Environnement), Université du Québec, 490 Rue de la Couronne, Québec, QC, Canada, G1K 9A9, Téléphone: (418) 261-6864, courriel: imene.ben_nasr@ete.inrs.ca

Fateh Chebana Professeur, Institut national de la recherche scientifique (Centre Eau Terre Environnement), Université du Québec, 490 Rue de la Couronne, Québec, QC, Canada, G1K 9A9, Téléphone: (418) 654-2542, courriel: Fateh.Chebana@ete.inrs.ca

A2. I. Ben Nasr and F. Chebana (2020). "Metaheuristic maximum pseudolikelihood method for the mixture copula models." Soumis au journal Stochastic Environmental Research and Risk Assessment.

Dans cet article, I. Ben Nasr a proposé une nouvelle approche d'estimation des copules mixtes multivariées. I. Ben Nasr a également mené une étude comparative entre l'approche proposée et une approche existante pour évaluer l'utilité de considérer ces deux approches dans un contexte hydrologique et identifier la meilleure. Le développement méthodologique et la rédaction ont été effectués par I. Ben Nasr. F. Chebana a fourni des commentaires durant l'exécution du travail et a révisé la version finale du manuscrit.

Abstract

Hydrological Frequency Analysis (*HFA*) has been widely applied to investigate the behavior and characteristics of hydrological variables especially regarding hydrological risk. Hydrological extreme events are characterized by several correlated random variables. For a better associated risk assessment, the dependence structure between these variables must be taken into account in a multivariate framework by considering copulas. On the other hand, hydrological variables might not respect the homogeneity condition in the usual multivariate framework where the events are generated from different phenomena. In such cases, the margins and/or copula may be affected. Hence, in particular, mixture copula should be considered. Recently, there have been an increasing number of studies dealing with the parameter estimation of mixture copula using the Expectation-Maximization method (*EM*). However, the *EM* algorithm has some drawbacks. To overcome these drawbacks, the main objective of this study is to propose a new parameter estimation approach for the mixture copula models, based on the maximum pseudo-likelihood using a meta-heuristic algorithm. A simulation study is conducted to evaluate the performance of the proposed method and to compare it with those of the *EM* method. It is found that the proposed method leads to the best results as compared to the *EM* method. Also, the results of Monte Carlo simulation for various sample sizes, demonstrate that the proposed method yields more accurate parameters estimates even with small sample sizes.

Keywords: Mixture copulas, estimation, Genetic Algorithm, Maximum pseudo-likelihood, simulations.

1. Introduction

Accurate estimation of quantile associated to extreme hydrologic events is of high importance since these events have significant risks to human beings and environment (e.g. Benameur *et al.*, 2017; Durocher *et al.*, 2016; Han *et al.*, 2014). Hydrological Frequency analysis (*HFA*) has been carried out widely for the prediction and risk assessment associated with hydrometeorological variables (e.g. Rao and Hamed, 2000). Traditionally, *HFA* is based on homogeneity as one of the fundamental assumptions, assuming that all observations are homogenous and thus belong to the same population (e.g. Buishand *et al.*, 2013; Willems, 2013; Yue *et al.*, 1999a). However, in a changing climate and due to human activities, such as river regulations and land use change, this assumption is not always fulfilled where the data series are heterogeneous (e.g. Das and Umamahesh, 2017; Jiang *et al.*, 2015; Sadegh *et al.*, 2015).

Heterogeneities in hydrological processes may be the result of a number of factors, including variations in regional weather patterns, natural and anthropogenic modifications of river systems, antecedent basin soil moisture (e.g. Evin *et al.*, 2011; Singh *et al.*, 2005b). Furthermore, a number of researchers found that hydrometeorological data are often the result of multiple sources or different physical generating mechanisms, such as different types of flood producing storms, rainfall and snowmelt floods, inundations and floodplain flow (e.g. Shin *et al.*, 2015; Smith *et al.*, 2011). Accordingly, analysis under (the non-fulfilled) homogeneity assumption cannot fully characterize the extreme events and may lead to inaccurate risk estimation. Therefore, to take into account the heterogeneity of hydrometeorological variables, the use of mixture distributions in *HFA* is suggested by many authors (e.g. Evin *et al.*, 2011; Shin *et al.*, 2016; Smith *et al.*, 2011; Yan *et al.*, 2016). All these studies have the common conclusion that mixture models can provide good fit to multi-source regimes and are essential for hydrology design under changing climate.

The aforementioned studies, and similar ones, even though highlighting the importance of considering the heterogeneity in *HFA* modeling, they focused on the univariate context, considering only one single variable of a given hydrological event such as annual flood peak or minimum flow. However, it has been showed over the last years that extreme hydrologic events can be characterized by the joint behaviour of several dependent variables (e.g. Chebana and Ouarda, 2011; De Michele *et al.*, 2013; Santhosh and Srinivas, 2013). Hence, univariate *HFA* does not procure a reliable assessment of the associated risk (e.g. Salvadori *et al.*, 2016; Volpi and Fiori, 2014). Therefore, the multivariate framework is becoming established in *HFA* studies.

In multivariate *HFA*, copulas is a key statistical tool and have been widely used in hydrology (e.g. Kao and Govindaraju, 2008; Requena *et al.*, 2013b; Zhang and Singh, 2006). To analyze data using a copula, its parameters should be estimated. For this purpose, several methods were proposed in the literature. Many authors support the use of the maximum pseudo-likelihood method (*MPL*) because it allows estimating the copula parameter independently from the marginal distributions. Furthermore, the *MPL* estimation method is much more generally applicable than the other methods since it does not require the dependence parameter to be real. Meanwhile, the main advantage of the *MPL* method is related to the small variance of parameters that the method provides that is why it is considered as the most efficient method. For a general review of copula theory and applications, the reader can refer to Joe (2014), Genest and Chebana (2017) and Singh and Zhang (2018). While copulas have been largely applied in *HFA*, their use in the lack of homogeneity is rarely studied, especially in *HFA* studies. However, at the best of our knowledge, the concept of heterogeneous multivariate hydrological series is not yet investigated in *HFA*.

In order to take into account the presence of the heterogeneity in the multivariate *HFA* modeling, mixture copula is an appropriate tool. They started to receive considerable attention and were

applied in several fields. For instance, in finance, Nguyen *et al.* (2016) and Hu (2006) examined complex dependencies between stock markets and gold prices using mixture copulas. Vrac *et al.* (2012) applied mixture copulas for climatic data clustering. Yu *et al.* (2013) modeled wind speed prediction using a Gaussian mixture copula. Despite all this diversity and progress in the mixture copula literature, little attention was devoted to this approach in water resources and *HFA*.

Proposed mixture copula models in the literature are limited to homogeneous mixture copula, *i.e.* the two components represent the same copula (particularly Gaussian copula) with different parameter values. It is based on the assumption that different sources of data have an identical dependence structure even though with different dependence strength (e.g. Bilgrau *et al.*, 2016; Yu *et al.*, 2013). However, hydrometeorological variables are often the results of multiple sources and it might be more rational to assume that individual sources follow different multivariate distribution (e.g. Fan *et al.*, 2016; Khan *et al.*, 2019). Heterogeneous mixture copulas (in which the two components represent two different copulas) are more general and more realistic. Their application in *HFA* is promising and may lead to a more accurate estimation of multivariate extreme events. Note that, in the univariate context, heterogeneous mixture distributions have been used with great success for the modeling of hydrometeorological variables (e.g. Calenda *et al.*, 2009; Hundedcha *et al.*, 2009; Shin *et al.*, 2016; Vrac and Naveau, 2007). Note that in some multivariate *HFA* studies, the terminology of mixture refers to the marginal distributions and not to the multivariate distribution or copula (e.g. Li *et al.* (2013), Khan *et al.* (2019) and Fan *et al.* (2016)).

In a preliminary step of a multivariate *HFA*, the homogeneity of the data should be tested. In cases where the data are homogenous, there would be no need to employ models that are more complicated. In the literature, different tests were proposed for multivariate homogeneity (see e.g. Ben Nasr and Chebana (2019a); Kojadinovic *et al.* (2016); and Quessy *et al.* (2013)). If the

homogeneity is rejected, the flood series can be considered as results of distinct flood-generating mechanisms or mixed populations. To analyze data using a mixture model, its parameters should be estimated. The estimation of the parameters can be performed through different methods. Therefore, the Expectation-Maximization (*EM*) algorithm has been used in estimating parameters of mixture models, in both univariate and multivariate contexts (e.g. Bilgrau *et al.*, 2016; Ding and Song, 2016; Dou *et al.*, 2016; Vrac *et al.*, 2012).

Despite its popularity, the *EM* algorithm suffers from several drawbacks such as divergence and poor accuracy, especially when dealing with small or moderate sample sizes (e.g. Ding and Song, 2016; Dou *et al.*, 2016). Furthermore, conventional methods may require several approximations, simplifications, or derivative information on functions of the model, and they may converge to local optimal solutions. Thus, there is a greater necessity to explore and apply new optimization methods to obtain optimal solutions. In the univariate case, the maximum likelihood method (*ML*), combined with a meta-heuristic algorithms for maximization (*MH*), is a good alternative to tackle limitations of the *EM* method (e.g. John, 2014; Song and Singh, 2010). During the past years, a variety of *MH* algorithms has been proposed in the literature, such as genetic algorithm (*GA*), particle swarm optimizations and harmony searches (e.g. Reca *et al.*, 2008). Among the existing *MH* algorithms, the *GA* has been widely applied in a variety of engineering applications such as in water resources (e.g. Joshi *et al.*, 2015; Karahan *et al.*, 2007; Reca *et al.*, 2008). The *GA* method is selected here based on its practical advantages. Indeed, *GA* algorithm requires only objective function values; no information about the gradient of the objective function is necessary. Despite strong evidence concerning the non-homogeneity of hydrological data, approaches dealing with mixture model estimation have not yet been considered in hydrology and specifically in a

multivariate context. In Table 3.1, the aforementioned estimation methods are summarised including their advantages and drawbacks.

Table 3.1. Overview of estimation method for parameter of mixture models

Model	Estimation method	Advantages	Drawbacks	Statistical and other fields	Hydrological field
Univariate mixture	Expectation-Maximization	-Easy to implement	-Need selection of the starting values	Arcidiacono and Jones (2003); Celeux <i>et al.</i> (2001); McLachlan and Jones (1988); Yuille <i>et al.</i> (1994)	Ailliot <i>et al.</i> (2009); Fan <i>et al.</i> (2016); Grego and Yates (2010); Yan <i>et al.</i> (2017); Zheng and Katz (2008)
		-Not requiring the computation of derivatives	-Slow convergence/divergence		
		-Numerically stable	-Convergence to local optima		
			-Not suitable for small samples		
	Maximum Likelihood	-More likely to reach the global optimum	-A fairly complex mathematical structure	Caudill and Acharya (1998); Everitt (2014); Leroux (1992); Van der Vaart (1996); Woodward <i>et al.</i> (1984)	Hundecha <i>et al.</i> (2009); John (2014); Leytham (1984); Li <i>et al.</i> (2012); Shin <i>et al.</i> (2015); Yan <i>et al.</i> (2019)
		-No initial values needed	-Computational effort		
-More efficient for small samples		-Huge derivation effort of information matrix			
Multivariate mixture copula	Expectation-Maximization	-Increasing likelihood in each iteration	-Same drawbacks as for the univariate mixture	Arakelian and Karlis (2014); Kim <i>et al.</i> (2013); Kosmidis and Karlis (2016); Qu and Lu (2019) séparer les ref une pr chaque drawb	the second aim of the present paper
		-Derivative-free optimizer	-Multiple local maximizers		
		-Dealing with parameters constraints implicitly	-Specifying marginal distributions		
			-Biased estimators		
			-Focus only on multivariate Gaussian distribution		
			-Focus only on Gaussian copula		
		Genetic Algorithm-Maximum Pseudo-likelihood	The specific aim of the present paper		

In the current study, for the multivariate setting, we propose a class of estimation methods for mixture copulas, which combine *MH* algorithms with maximum pseudo-likelihood (*MH-MPL*). Given the advantages of *GA* as a particular case of *MH*, the proposed *MH-MPL* incorporates the *GA* and the well-known *MPL* estimation method. In the univariate context, there are few studies considering *GA* to estimate parameters of flood frequency distributions (e.g. Shin *et al.*, 2014; Shin *et al.*, 2015). However, despite its popularity, *GA* has not yet been used in estimating parameters in multivariate mixture models and in particular in *HFA* applications.

The remainder of this paper is organized as follows. Section 2 briefly reviews the theoretical background on mixture copulas. Section 3 deals with parameter estimation methods in mixture copulas. Then, the simulation study is presented in section 4. Conclusions are given in section 5.

2. Copula and mixture of copulas

This section provides an overview of the main concepts about copula and mixture models

2.1. Copula model

Copulas are proposed as a flexible tool for constructing multivariate distributions and modeling the dependence structure between correlated variables and widely used in hydrology (e.g. Genest and Chebana, 2017; Hao and Singh, 2016). The theory of copulas is based on the Sklar's theorem (Sklar, 1959) which in the case of a bivariate case can be represented as:

$$H(x, y) = C(F(x), G(y)), x, y \in \mathfrak{R} \quad (3.1)$$

where $H(x, y)$ is the joint cumulative distribution function of the random variables X and Y , $F(x)$ and $G(y)$ are the marginal distribution functions of X and Y , respectively, and the mapping function $C: [0,1]^2 \rightarrow [0,1]$ is the copula function. Accordingly, the density function of the copula C can be defined as

$$c(u, v) = \frac{\partial^2 C(u, v)}{\partial u \partial v} \quad (3.2)$$

where (u, v) are pseudo-observations such that $u_i = \frac{R_i}{n+1}$; $v_i = \frac{S_i}{n+1}$ and R_i is the rank of X_i^1 among $X_1^1, \dots, X_n^1$; S_i is the rank of the concomitant X_i^2 among $X_1^2, \dots, X_n^2$.

A large variety of families of copulas are available to model multivariate dependencies. Overall, there are two families of copulas that are widely applied in *HFA*: the extreme value copulas, and the Archimedean copulas (e.g. Chebana and Ouarda, 2009; Requena *et al.*, 2013b; Salvadori and De Michele, 2010). Indeed, these copulas present several desirable properties. The popularity of the Archimedean family stems from several desirable properties: 1) can be easily generated, 2) are symmetric and associative, and 3) include a large variety of copulas. Clayton and Frank copulas are the most used ones from the Archimedean family to characterise the dependence structure between hydrological variables. Meanwhile, given their importance in modeling catastrophic events, extreme value (*EV*) copulas have been extensively used in recent years (e.g. Kao and Govindaraju, 2008; Salvadori and Michele, 2011). A large number of studies considered the Gumbel copula as the *EV* copula that best represents the relation between hydrological variables (e.g. Lee and Salas, 2011; Zhang and Singh, 2006).

2.2. Mixture copula models

The modeling of a mixture of distributions has long been restricted upon Gaussian distributions and it is only recently that researchers have started to study broader general models (e.g. Khan *et al.*, 2019; Yu *et al.*, 2013). Among these extensions, models featuring copulas are still rare, but certainly promising. For the sake of completeness, since the focus is on the dependence structure, the basic definitions and interpretations of the multivariate mixture are briefly described as follows. Let (X, Y) be a bivariate random vector with a sample of size n from a finite mixture model with

two components. The bivariate cumulative distribution function (*CDF*) can be expressed as (e.g. Qu and Lu, 2019; Vrac *et al.*, 2005).

$$H(x, y) = \omega H_1(x, y, \theta_1) + (1 - \omega) H_2(x, y, \theta_2), \omega \in [0, 1] \quad (3.3)$$

Therefore, from Sklar's theorem in equation (3.1), there exists a copula C such that

$$H(x, y) = \omega * C_1((F_1(x), G_1(y), \theta_1) + (1 - \omega) * C_2(F_2(x), G_2(y), \theta_2), \omega \in [0, 1] \quad (3.4)$$

where ω is the mixture ratio, F_k and G_k are univariate marginal distributions of the univariate data X and Y and θ_k is the parameter of the copula corresponding to the k^{th} component.

We note that $u_i := \frac{R_i}{n+1} \sim \text{Unif}(0,1)$ and $v_i := \frac{S_i}{n+1} \sim \text{Unif}(0,1)$ for $i = 1 \dots n$; such that

R_i is the rank of x_i among $x_1, \dots, x_n$ and S_i is the rank of y_i among $y_1, \dots, y_n$ (e.g. Nelsen, 2013).

Then, when it exists, the probability density function (*pdf*) h of the mixture model is defined as:

$$h(x, y) = \omega * c_1(u_1, v_1, \theta_1) * f_1(x) * g_1(y) + (1 - \omega) * c_2(u_1, v_1, \theta_2) * f_2(x) * g_2(y) \quad (3.5)$$

where c_k is the pdf of the copula C_k , θ_k is the parameter of the copula and f_k, g_k are the marginal pdfs. Note that all the above densities exist since usually in *HFA*, the variables to study are continuous. Consequently, following Thongkairat *et al.* (2019) and Vrac *et al.* (2012), the two-components mixture copula, as well as the corresponding *pdf*, are given by:

$$C_{mix}(u, v, \omega, \theta) = \omega * C_1(u, v, \theta_1) + (1 - \omega) * C_2(u, v, \theta_2) \quad (3.6)$$

$$c_{mix}(u, v, \omega, \theta) = \frac{\partial^2 C_{mix}}{\partial u \partial v} = \omega * c_1(u, v, \theta_1) + (1 - \omega) * c_2(u, v, \theta_2) \quad (3.7)$$

The assumption used in conventional mixture models is that the two copulas C_1 and C_2 belong to the same family (e.g. Bilgrau *et al.*, 2016; Bonanomi *et al.*, 2019; Yu *et al.*, 2013). In this case, the model is referred as homogenous mixture. However, under heterogeneity assumption, not only the statistical parameters but also the type of copula could be changing. Hence, single copula is unable

to convey the hydrological behaviour comprehensively, especially regarding the tail dependence. This complex behaviour suggests that the data can be described as a realization of a heterogeneous mixture, allowing different shapes for C_1 and C_2 . Despite previous research efforts regarding multivariate mixture models, it is important to mention that heterogeneous multivariate mixture models have been rarely considered (e.g. Christensen *et al.*, 2019; Vrac *et al.*, 2012) and not yet been considered in *HFA* context.

Since the focus is mainly on extreme events, the main advantage of using heterogeneous mixture copulas is that we can create a copula with different characteristics, especially in the upper and lower the tails. For example, if the two copulas in the mixture have opposite tail dependence structures, such that the Clayton and Gumbel copulas, the resulting mixture copula accounts for both upper and lower tail dependencies. Supporting this point, Bonanomi *et al.* (2019) and Christensen *et al.* (2019) showed that the mixture copula inherits characteristics from its component copulas. Particularly, the upper and lower tail dependence coefficients can be estimated by (e.g. Bonanomi *et al.*, 2019)

$$\lambda_U^{mix} = \omega \lambda_U^{C_1} + (1 - \omega) \lambda_U^{C_2} \quad (3.8)$$

$$\lambda_L^{mix} = \omega \lambda_L^{C_1} + (1 - \omega) \lambda_L^{C_2} \quad (3.9)$$

where $\lambda_U^{C_k}$ and $\lambda_L^{C_k}$ are the upper and lower tail dependence coefficients of the copula C_k , $k = 1, 2$.

3. Parameter estimation methods

The estimation of the parameters of the mixture model, that best fits the data, can be performed through different methods. In the following, a brief description of the used methods in the present study is presented.

3.1. Expectation-Maximization (*EM*) method

The *EM* algorithm is an iterative method, initially conceived to compute the maximum likelihood estimates of parameters of a model when some observations are missing. The *EM* algorithm has dominated the literature on maximum likelihood estimation of mixture models, in both univariate and multivariate setting (e.g. Ding and Song, 2016; Dou *et al.*, 2016; Hu, 2006; Shin *et al.*, 2015; Yan *et al.*, 2016). Throughout this paper, since the focus is on the dependence structure, the following *EM* algorithm is applied to maximize the pseudo-likelihood function instead of the complete likelihood, in which the empirical marginal distributions are used instead of the parametric marginal distributions (e.g. Genest *et al.*, 1995). Hence, the marginal distributions are estimated as $F_n(x) = \frac{1}{n+1} \sum_{i=1}^n 1_{(X_i \leq x)}$ and $G_n(y) = \frac{1}{n+1} \sum_{i=1}^n 1_{(Y_i \leq y)}$. Therefore, the corresponding density functions f and g are estimated using kernel density estimators (e.g. Adamowski, 1985; Lall, 1995; Santhosh and Srinivas, 2013).

The *EM* algorithm consists mainly in two steps. Firstly, in the *E*-step, the posterior probability of the i^{th} observation belonging to the k^{th} component is estimated by

$$E\text{-step: } \pi_{ik}^{(l+1)} = \frac{\omega_k^{(l)} f_k(x_i) g_k(y_i) c_k(u_i, v_i, \theta_k^{(l)})}{\sum_{k=1}^2 \omega_k^{(l)} f_k(x_i) g_k(y_i) c_k(u_i, v_i, \theta_k^{(l)})}, \quad i = 1, \dots, n, \quad k = 1, 2 \quad (3.10)$$

Then, in the *M*-steps, the parameter estimates are updated using the estimated probabilities, obtained in the *E*-step, as given by

$$M\text{-step 1: Set } \omega_k^{(l+1)} = \frac{1}{n} \sum_{i=1}^n \pi_{ik}^{(l+1)} \quad (3.11)$$

$$M\text{-step 2: Update } \theta_k^{(l+1)} = \underset{\theta}{\operatorname{argmax}} \left(\sum_{i=1}^n \log \left(\sum_{k=1}^2 \pi_{ik}^{(l+1)} c_k(u_i, v_i, \theta_k^{(l)}) \right) \right) \quad (3.12)$$

The algorithm iterates between the E -step and the M -steps until some convergence criterion is satisfied. In this paper, the algorithm is considered to converge when the change of parameter values is less than a predefined small threshold value.

3.2. Meta-heuristic Maximum Pseudo-Likelihood method

Although the EM algorithm is a popular method to estimate parameters of mixture (univariate and multivariate) distribution, this method has some limitations (see Table 3.1). Especially, it does not converge when dealing with small-sized samples (e.g. Ding and Song, 2016; Liu *et al.*, 2019), which is typically encountered in hydrology. Furthermore, the convergence of the EM algorithm to the global maximizer depends on the starting point of the algorithm. To overcome these drawbacks, we propose the (MPL) method combined with a MH algorithm for the optimization of the parameters (noted as $MH-MPL$). The mixed copula is fitted to data by maximizing the logarithmic pseudo-likelihood function, defined by

$$LL(\omega, \theta) = \sum_{i=1}^n \log \sum_{k=1}^2 \omega_k c_k(u_i, v_i, \theta_k) \quad (3.13)$$

For the reasons presented above, in the current study, a GA was used as a representative of MH algorithms. This choice is also motivated by the presence of a significant number of publications that indicated that the GA provides a good performance in the estimation of the parameters of non-mixture and mixture univariate distributions (e.g. Hassanzadeh *et al.*, 2011). A GA consists of genes, individuals, populations, and generations. In the current study, the population size and the maximum generation number in the GA are taken equal to 500, which is considered large enough to obtain robust estimators (e.g. Song and Singh, 2010). GA starts with a random population of trial solutions. The fitness value (called objective function) associated with each chromosome is evaluated. In our case, this function corresponds to (3.13). Then the new population for the next

generation is obtained using three genetic operations, namely selection, crossover and mutation. More details on *GA* and its procedure can be found in Sivanandam and Deepa (2008).

The proposed *MH-MPL* method has several advantages. Indeed, it can reach the global optimum without requiring initial values for the parameters of the mixture copulas (e.g. Song and Singh, 2010). Furthermore, this approach is efficient in estimating parameters for small sized samples. Moreover, as supported by Reça *et al.* (2008), *MH* algorithm does not require derivatives of the objective function and can hence be applied to solve complex and discontinuous optimization problems. Therefore, it can be appropriate to various kinds of mixture model with different component copulas including situations where derivative of (3.13) does not have explicit form. Further, the *MH-MPL* method can overcome the problem of trapping at local optima, which is common in some classical gradient-based methods. It is also interesting to mention that, due to its considerable flexibility, the proposed *MH-MPL* method can be considered also for the estimation of non-mixture copula. For instance, it has been successfully applied in the estimation of one-parameter copula. Further, Reddy and Singh (2014) showed that the *MH-MPL* is more accurate than classical methods for estimating the parameter of the copula.

4. Evaluation of estimation methods via simulation

The aim of the simulation study is twofold. First, we evaluate the performance of the proposed estimation method in the hydrological context. Second, we compare the performance of the *MH-MPL* and *EM* methods. To this end, we consider practical situations commonly encountered in hydrological applications.

4.1. Simulations design

A Monte Carlo simulation was conducted to evaluate the performance of a parameter estimation method by generating and analysing samples from various models with known parameters. Generally, a flood event is characterized by three main features: peak (Q), volume (V) and duration (D). It was pointed out in the available literature on flood events that generally Q and D are not significantly dependent whereas the most correlated variables are Q and V (e.g. Aissia et al., 2012; Fu and Butler, 2014). Hence, these two variables are considered in this simulation study. Three different copulas, commonly used in *HFA*, are used to model the dependence structure between Q and V , namely, Clayton, Frank and Gumbel copulas. The mixing ratios $\omega=0.2, 0.3, 0.5$ are considered for building the mixture copula model. It is worth noting that these copulas are considered under different scenarios to generate heterogeneous data. Hereafter, the considered scenarios are:

- a) Homogenous mixture model: Clayton-Clayton ($C-C$), Frank-Frank ($F-F$) and Gumbel-Gumbel ($G-G$);
- b) Heterogeneous mixture model: Clayton-Frank ($C-F$), Clayton-Gumbel ($C-G$), Frank-Gumbel ($F-G$).

Since the dependence between two variables is described by both the dependence type (copula family) and the dependence strength, different values of Kendall's τ are used in this study. Therefore, three values of $\tau=0.2, 0.6, 0.8$, are considered corresponding to weak, moderate, and strong dependence, respectively. These values are selected on the basis of situations commonly encountered in *HFA* (e.g. Requena *et al.*, 2013b; Zhang and Singh, 2007b). The true parameter of each component copula is defined to match the corresponding range of dependence and is estimated using Kendall's tau inversion method (e.g. Nelsen, 2013).

Besides the dependence strength, sample size is a relevant factor for the performance of a parameter estimation method. Hence, a sensitivity analysis of the performance of aforementioned methods is performed regarding the sample size. Since hydrological series are typically characterised by small sized samples, the assessment of the behaviour of each method was performed under $n = 30, 50, 100$. The values of n are selected on the basis of cases frequently occurred in practical situations (see for examples series in Barth *et al.* (2017) and Santhosh and Srinivas (2013)). For each combination of mixture copula, sample size and Kendall' τ , we generate $M=1000$ synthetic series through Monte Carlo simulations. To generate synthetic data from a given mixture copula, we applied the procedure proposed by Nelsen (2013), based on the conditional distribution method. A diagram of the simulation study is shown in Figure 3.1.

In order to evaluate the performance of each estimation method and to compare them, the relative error (RE), relative root-mean square error ($RRMSE$) and the relative $BIAS$ ($RBIAS$) are calculated from the true parameters and the estimated ones as

$$RE_j = \frac{\hat{\theta}_j^i - \theta_j}{\theta_j} * 100 \quad (3.14)$$

$$RRMSE_j = 100 \sqrt{\frac{1}{M} \sum_{i=1}^M \left(\frac{\hat{\theta}_j^i - \theta_j}{\theta_j} \right)^2} \quad (3.15)$$

$$RBIAS_j = \frac{1}{M} \sum_{i=1}^M \left(\frac{\hat{\theta}_j^i - \theta_j}{\theta_j} \right) * 100 \quad (3.16)$$

where j is the number of parameters, M is the number of Monte-Carlo samples, θ_j is the true parameter and $\hat{\theta}_j^i$ is the estimated parameter from sample i of each simulation.

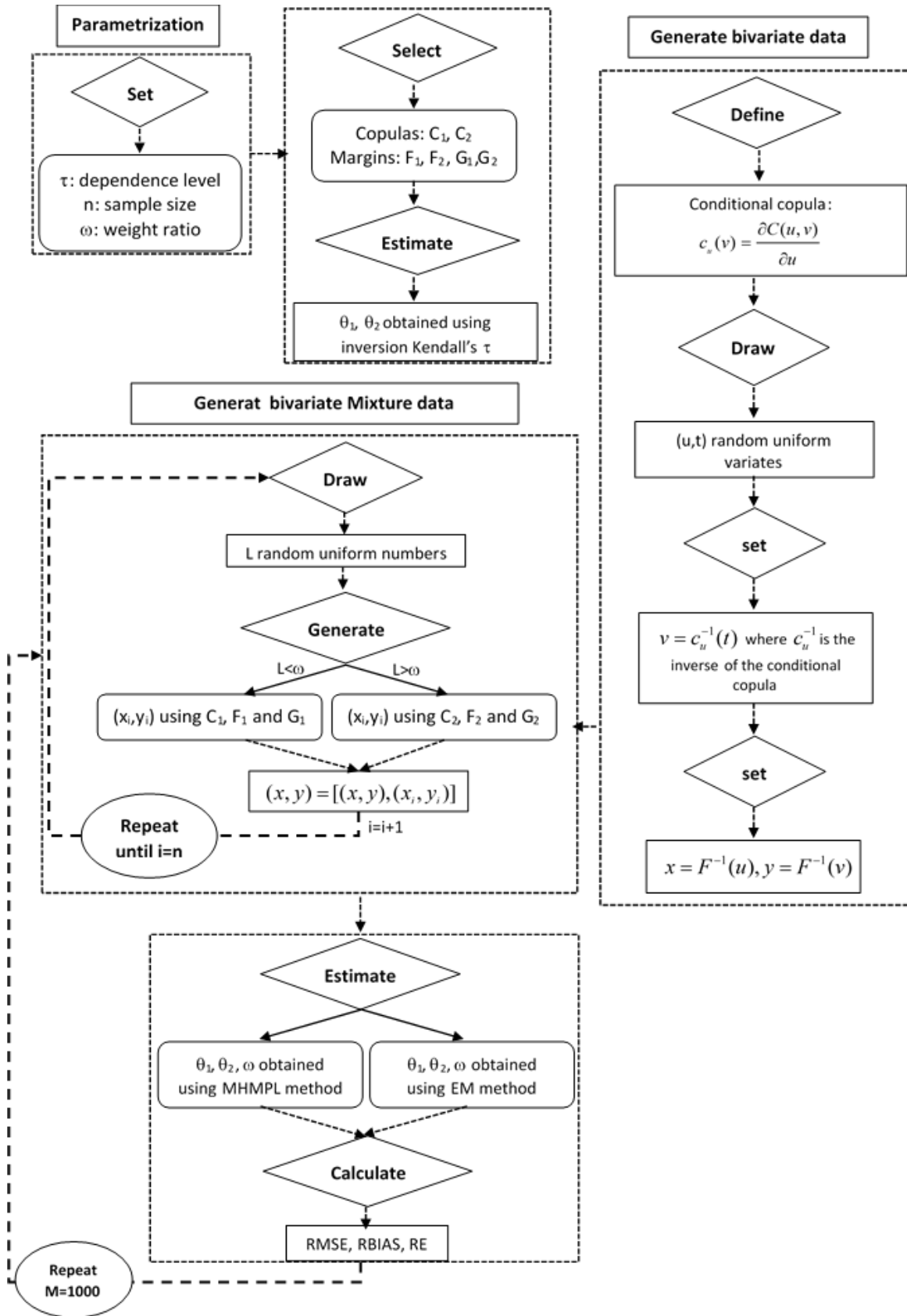


Figure 3.1. Flowchart of the simulations study

4.2. Simulation results

To compare the performance of the two methods, they were applied to synthetic series. Then, different performance criteria were computed. The first step of the assessment of the performance of an estimator is the analysis of the effect of the sample size on its behaviour. The corresponding results are presented in Table 3.2. From this table, one can see that overall *RRMSE* and *RBIAS* decrease when the sample size n increases, for both *EM* and *MH-MPL* approaches. As an example, when estimating the weight ratio ω in the case of the Frank-Gumbel mixture copula, for the *MH-MPL* estimators, the *RRMSE* and *RBIAS* decrease from 31.3 to 11% and from 12.9 to 4.7%, respectively, when the sample size n increases from 30 to 100 and for a dependence level $(\tau_1, \tau_2) = (0.8, 0.8)$. The same results hold true when estimating copulas' parameter θ_1 and θ_2 (Table 3.2). These findings are consistent with other studies dealing with copula parameter estimation (e.g. Genest *et al.*, 1995).

From Table 3.2, it is also found that *EM* estimators show larger variability in terms of *RBIAS* and *RRMSE*, especially when $n < 100$. The *EM* method seems to suffer from large biases and large *RRMSEs* for smaller samples sizes yielding extremely inaccurate estimates. For example, when estimating the weight ratio ω of the Clayton-Frank model for $n=30$, the *RRMSEs* are 67.3 and 30.5%, for the *EM* and *MH-MPL* methods, respectively. Furthermore, the *EM* estimation method does not behave correctly neither for $n=30$ nor $n=50$. Indeed, in these cases, the *EM* method leads to a large underestimation (negative *RBIAS*) of the parameters of the mixture model. For instance, the Clayton-Frank model the *RBIAS*, associated to the *EM* estimates of the copulas' parameter θ_1 and θ_2 , are equal to -27.4 and -36.9%, when $n=30$. However, the *MH-MPL* estimates appear to suffer slightly less under reduced sample size than the *EM* ones. Indeed, for the six considered mixture models, the *RBIAS* associated to the *MH-MPL* method is smaller by a factor of two or

more than the *RBIAS* of the *EM* estimator. This agrees with similar findings by John (2014), when dealing with univariate mixture models. In the same vein, Kim *et al.* (2013) and Arakelian and Karlis (2014) provided also simulation results of the performance of the *EM* method in the context of mixture copula models. Based on their results, they reached the same conclusion that the *EM* method is not recommended for situations with small sample sizes ($n < 150$).

In order to provide a visual support of the behaviour of each estimation methods, boxplots of *RE* are presented in Figure 3.2. As expected, the minimum sample size required to have a reliable estimation in the case of the *MH-MPL* method is less than that needed for the *EM* method. Actually, the *MH-MPL* method always has the smallest *RE* while the *EM* method always has the largest *RE*, regardless of the sample size. These findings hold for the estimation of the weight ratio ω as well as the copulas' parameter θ_1 and θ_2 . In all cases, both the median value of *EM* estimators and its variability represented by the height of the box are the largest. In conclusion, the *MH-MPL* method is more accurate in detecting the characteristics of the mixture model than the *EM* method, regardless of sample size and mixture copulas.

Besides the sample size, the dependence strength (as described by Kendall's τ) is a relevant factor to the performance of an estimation method. Results of the effect of varying Kendall's τ are presented in Table 3.3. In order to avoid the interference of the sample size n and the weight ratio ω in the assessment of the performance of different estimation methods, these factors are taken to be equal to $n=150$ and $\omega=0.5$, respectively. It can be seen from Table 3.3 that, for a given method, the estimation of the weight ratio ω is not affected by the dependence level of each mixture component. Indeed, both estimation methods provide similar results, regarding the *RBIAS* and *RRMSE*, when Kendall's τ increases from (0.2, 0.6) to (0.6, 0.8). For instance, in the case of the

MH-MPL method, the *RBIAS* and *RRMSE* are around 1.8% and 7.3% for the Frank-Gumbel mixture, regardless of the dependence level.

Table 3.2. RRMSE (%) and RBIAS (%) of different estimators, over 1000 replications, for different mixture models and different sample size when dependence level $\tau = (0.2, 0.6)$ and weight ratio $\omega = 0.5$

Copula		EM method						MH-MPL method					
		30		50		100		30		50		100	
		RRMSE	RBIAS	RRMSE	RBIAS	RRMSE	RBIAS	RRMSE	RBIAS	RRMSE	RBIAS	RRMSE	RBIAS
ω	C-C	64.7	-37.2	44.8	-26.9	26.4	-8.3	33.1	17.9	27.8	8.4	19.1	5.9
	F-F	65.2	-38.1	46.4	-30.5	25.6	-7	31.1	10.3	24.7	7.9	15.3	5.1
	G-G	70.2	-28	43.1	-28.1	26.6	8.4	34.1	14.5	23.8	7.8	14.7	5.2
	C-F	67.3	-49.1	41.9	-26.8	20.2	-13.1	30.5	13	23.6	7.9	10.2	4.4
	C-G	64.1	26.4	43.6	19.7	25.8	11.9	30	10.9	25.8	6.2	12.5	3.3
	F-G	64.5	-28.7	42.1	20	24.9	-10.1	31.3	12.9	23.9	7.4	11	4.7
Θ_1	C-C	75.6	-22.5	44.9	-11.3	22.8	6.1	33.4	-13.6	28.7	-9	12.3	4.6
	F-F	74.2	-25.4	45.5	-10.7	23.3	-5.4	34.1	-16.8	31.8	-8.9	13.4	5.6
	G-G	71.2	-27.4	42.4	-16.9	20.1	6.9	34.6	-14.3	30.4	-9	12.5	5.2
	C-F	70.3	-22.8	44.7	-15.1	23.5	-9.3	33.4	13.4	30.5	8.7	12.1	3.4
	C-G	71.1	-25.8	45.9	-17.5	22.4	-9.6	34.7	14.1	28.6	9.4	10.6	3.9
	F-G	73.2	-28.2	40.1	-12.3	21.3	-8.5	32.5	11.2	29.4	7.4	12.3	4.4
Θ_2	C-C	71.5	-25.8	43.1	-17.7	23.8	6.5	32.1	-15.3	28.3	-7.9	12.4	-5.1
	F-F	72.9	23.9	46.4	11.9	22.4	5.6	33.3	-17.8	29.9	-9.6	11.5	-5.2
	G-G	71.4	-36.9	43.6	-20.9	23.1	-5.4	36.6	12.9	27.5	-7	13.2	-5.6
	C-F	73.2	-21.8	42.9	-14.8	20.8	-8.2	33.9	10.9	26.4	7.1	15.3	4.4
	C-G	72.1	53.9	45.1	28.6	23.9	-8.1	32.2	13.6	22.3	8.4	11	4.7
	F-G	71.3	24.1	41.1	17.2	24.1	-8.5	33.4	17.3	26.4	8.8	10.8	3.6

This table presents the RBIAS and RRMSE for different mixture of copula model for different sample sizes, using the metaheuristic maximum pseudo-likelihood (MH-MPL) and EM methods. Light gray color corresponds to best results regarding RBIAS values and dark gray color to best results regarding RRMSE.

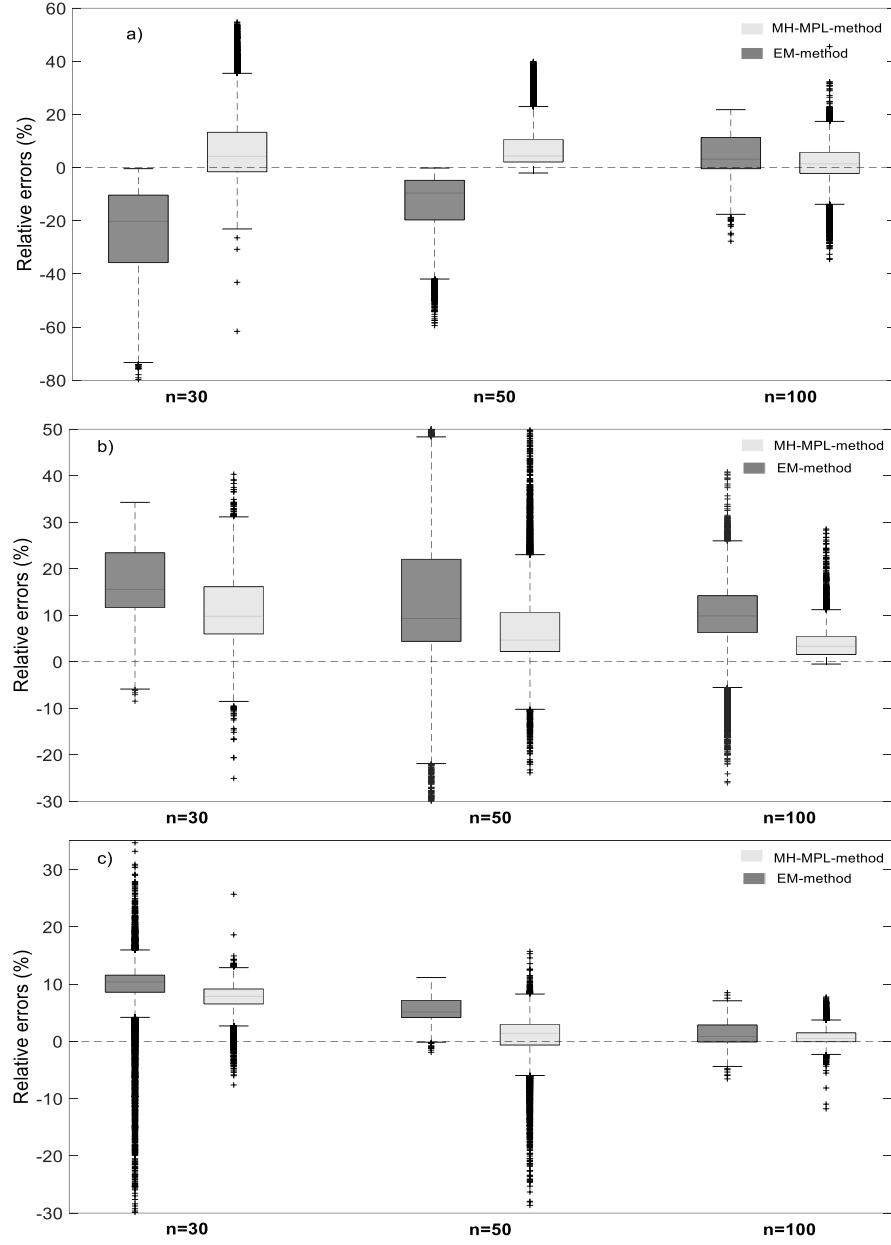


Figure 3.2. Relative errors results of MH-MPL and EM estimators for different sample size n , dependence level $\tau = (0.6, 0.8)$ and $\omega=0.5$: a) weight ratio ω , b) first copula parameter θ_1 and c) second copula parameter θ_2

Table 3.3. *RRMSE* (%) and *RBIAS* (%) of different estimators, over 1000 replications, for different mixture models and different dependence level when sample size $n=150$ and weight ratio $\omega=0.5$

Copula		<i>EM method</i>						<i>MH-MPL method</i>					
		(0.2, 0.6)		(0.2, 0.8)		(0.6, 0.8)		(0.2, 0.6)		(0.2, 0.8)		(0.6, 0.8)	
		RRMSE	RBIAS	RRMSE	RBIAS	RRMSE	RBIAS	RRMSE	RBIAS	RRMSE	RBIAS	RRMSE	RBIAS
ω	C-C	17.9	-4.2	17	-4.2	17.4	-3.7	7.9	2.8	7.5	2.1	7.5	2
	F-F	16.2	-3.3	15	-3.3	14.7	-3.3	8.1	2.5	7.7	2.4	7.2	2.3
	G-G	16.8	-4.5	16.7	-4.3	16.0	-4.2	8.2	1.9	8.5	1.6	7.9	1.8
	C-F	19	-4.5	18.5	-4.2	18.8	-3.8	7.6	1.7	7.6	1.6	7	1.4
	C-G	18.7	-5.9	18.7	-4.5	18.5	-4.8	8.2	1.2	8.1	1.1	7.4	1
	F-G	19.9	-5.8	19.4	-5	19.2	-4.7	7.7	1.9	7.3	1.8	7	1.8
Θ_1	C-C	33	-11.2	30.6	-10.2	20.5	-4.2	29.4	-7.6	29.1	-5.2	9.4	2.8
	F-F	34.1	-11.3	30.6	-11.5	20.7	7.3	29.6	-8.1	29.8	-7.1	8.5	3.1
	G-G	33.9	-10.5	30.7	-10.2	22.2	5.7	29.1	-9.6	28.4	-6.3	8.8	3.2
	C-F	30.6	-16.3	31.2	-13.8	20.5	-5.4	27.7	-5.7	27.1	-4.5	8.3	2
	C-G	31.2	-14.3	31.1	-13.2	22.7	-6.4	28.2	-7.2	28.7	-6.4	7.8	1.9
	F-G	30.3	-12.6	34.7	-12.4	21.4	-4.8	27.1	-4.9	27	-5	8.3	3.2
Θ_2	C-C	25.1	-6.36	22.4	-4.3	20	-4.2	18	3.1	9.7	2.8	9.2	2.5
	F-F	23.5	7.01	20.7	4.2	19.5	4.2	17.6	2.7	9.8	2.3	9.4	2.4
	G-G	24.3	6.35	20.5	4.5	20.2	3.9	18.7	3.3	8.8	1.6	8.5	1.7
	C-F	24.2	-5.72	20.8	-4.1	22.8	-3.5	16.4	3.2	9.1	2.4	7.6	2.2
	C-G	23.8	9.54	22.3	4.4	23.4	3.7	17.4	4.1	7.9	2.2	7.3	2.3
	F-G	24.1	8.48	22.1	3.9	21.3	3.9	18.5	3.4	9.2	1.7	8.3	1.8

This table presents the *RBIAS* and *RRMSE* for different mixture of copula model with different dependence strength, using the metaheuristic maximum pseudo-likelihood (*MH-MPL*) and *EM* methods. Light gray color corresponds to best results regarding *RBIAS* values and dark gray color to best results regarding *RRMSE*.

As can be seen from Table 3.3, when estimating the weight ratio ω for homogenous mixture copulas, there is a noteworthy difference in *RBIAS* criteria between the two estimation methods, regardless of the dependence strength, as described by Kendall's τ . Furthermore, this difference increases to about 65% for estimating the heterogeneous mixture copulas. Moreover, the same trend is observed concerning the difference in *RRMSE*. Once again, such a difference becomes slightly larger for heterogeneous mixture copulas. Accordingly, the *MH-MPL* method consistently outperforms the *EM* method in estimating the weight ratio. These results agree with the findings of other studies dealing with univariate mixture models, which found a significant improvement in the performance of the *MH-MPL* method as compared with the *EM* algorithm (e.g. John, 2014).

Results from Table 3.3 highlight also the effect of the dependence level on the estimation of the copulas' parameter. Overall, the performance of the two approaches, regarding the estimation of copulas' parameter θ_1 and θ_2 , is better when the dependence level is stronger. For instance, the *RBIAS* and *RRMSE* are larger for small dependence level, as expected. For example, in the case of the Clayton-Gumbel mixture, the absolute value of the *RBIAS* of the *MH-MPL* estimator decreases from 7.2 to 1.9% for the first parameter θ_1 and from 3.4 to 1.8% for the second parameter θ_2 , when the dependence level increases from $(\tau_1, \tau_2) = (0.2, 0.6)$ to $(\tau_1, \tau_2) = (0.6, 0.8)$. This finding is also valid for the *RRMSE* where a smaller value of dependence level is associated to larger *RRMSE*. These results are consistent with those of other studies dealing with estimation of non-mixture copula parameters such as Genest *et al.* (1995) and Brahimi and Necir (2012) who showed that an estimation method has a better performance for stronger dependence level.

Additionally, when estimating θ_1 for homogenous mixture copulas, the difference in *RBIAS* between *MH-MPL* and *EM* method is about 25%, for $(\tau_1, \tau_2) = (0.2, 0.6)$. Again, this difference increases as the dependence strength increases to reach 50% when $(\tau_1, \tau_2) = (0.6, 0.8)$. These results support the robustness of the *MH-MPL* when estimating copulas' parameter. Its performance is not highly affected even for low dependence level. This finding is in agreement with the findings in Kim *et al.* (2007) which showed that, in the case of non mixture copula, the *MPL* method is more suitable for the estimation of the copula's parameter with low dependence strength. Moreover, the relative *RMSE* seem to be considerably higher for the *EM* method, when estimating copulas' parameter. Meanwhile, previous discussion about comparison of *EM* and *MH-MPL* results regarding the *RBIAS* also holds for the *RRMSE*. These results are likely to be related to the fact that the convergence of the *EM* algorithm to the global maximum depends on the starting point of the

algorithm. This might be also explained by the fact that the two component copulas in the mixture model get further away from each other when the Kendall's τ increase.

Since both *EM* and *MH-MPL* performance criteria improve with increasing dependence strength and in order to identify the best-recommended one for hydrological applications, a deeper analysis is carried out to compare the two approaches. For this purpose, in order to show complete information and to present an overview of the results, boxplots of the effect of Kendall's τ on the relative errors are provided in Figure 3.3. Several conclusions can be drawn from these boxplots. From Figure 3.3a), it is found that the estimate of ω is not very sensitive to the dependence level. In fact, for the *MH-MPL* method, the median values (line inside the box) is almost the same, for the three dependence levels combinations. The same holds for *EM* method, regardless of the dependence level. It is also found that, in comparison to the *EM* method, *MH-MPL* method has an overall better behavior. Indeed, *EM* estimators show higher median value. This highlights the low ability of the *EM* method to correctly estimate the weight ratio ω in certain circumstances. Note that although the median is the value considered to assess the performance of each method, the variability in the results (i.e., the height of the boxes) should also be considered in the decision process, as it refers to the uncertainty in the results given by the estimation method. Taking into account all the information provided by the aforementioned criteria, the *MH-MPL* method was identified as the best in terms of estimating the weight ratio, as both the median value and its variability are the smallest.

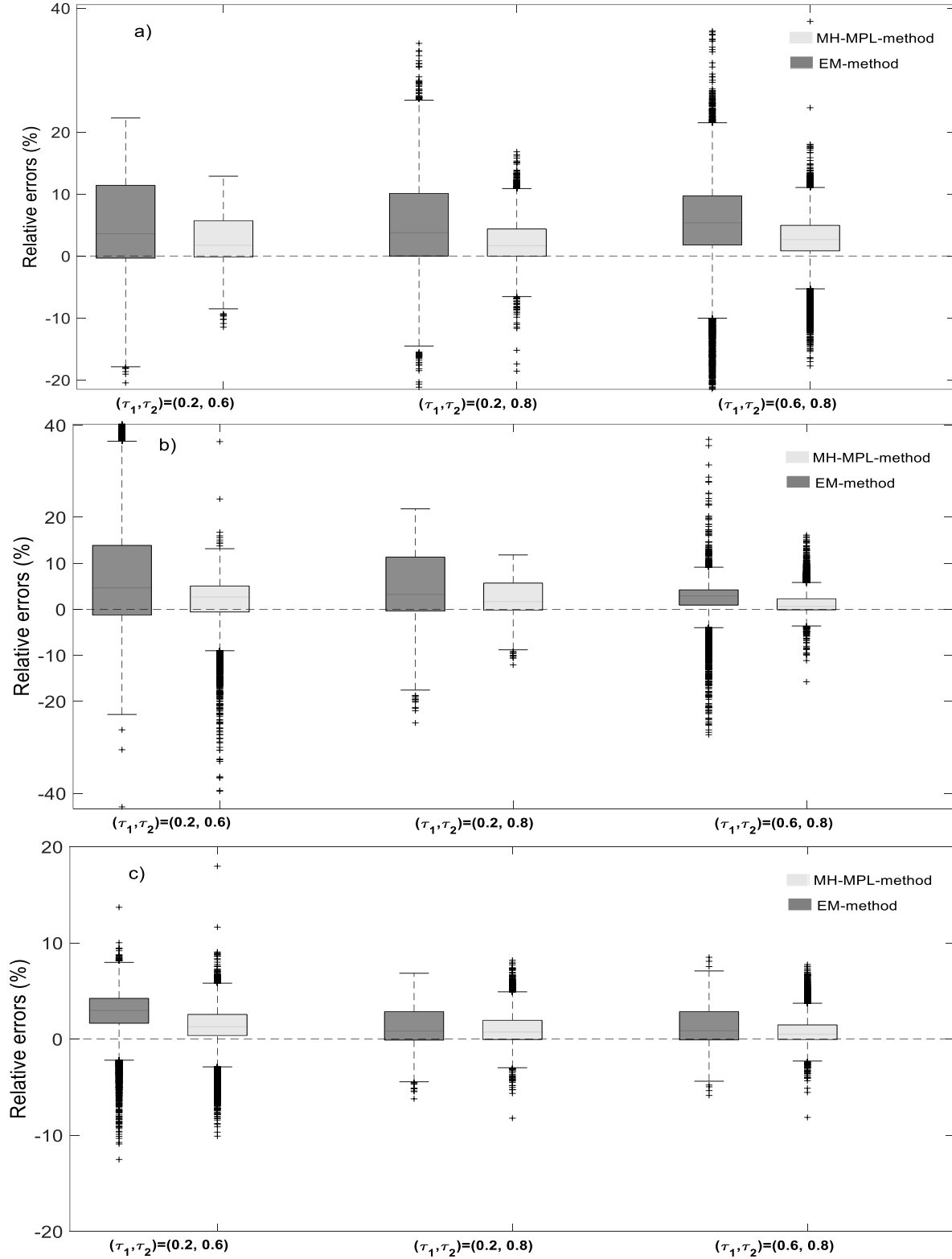


Figure 3.3. Relative errors results between MH-MPL and EM estimators for different kendall's τ , $n=150$ and $\omega=0.5$: a) weight ratio ω , b) first copula parameter θ_1 and c) second copula parameter θ_2

We can see from Figures 3.3b) and 3.3c) that, as expected, as the dependence gets stronger, the *RE* decrease for the two estimation methods. This observed finding mirrors those of the previous studies that have examined different estimation method for non-mixture copulas (e.g. Brahimi and Necir, 2012; Capéraà *et al.*, 1997). However, although the performance of parameter estimates from the two methods show similar patterns, the *MH-MPL* method leads to better results when estimating copulas' parameter. Indeed, regardless the dependence level, the reduction in the *RE* of the *MH-MPL* method is remarkable. The *RE* from the *MH-MPL* method is about 20% smaller than that from the *EM* method. Beyond confirming the obvious dominance of the *MH-MPL* estimator and the poor performance of the *EM* estimator, the results in Figure 3.3b) and 3.3c) indicate that the *EM* method leads to larger uncertainty especially for low dependent data ($\tau=0.2$).

The mixing proportion ω is of primary interest. Results regarding the effect of ω on the performance of each estimation method are displayed in Table 3.4.

To discard the effect of the sample size and the dependence level, we analyse the effect of ω when $n=150$ and $(\tau_1, \tau_2) = (0.6, 0.8)$. As can be seen from Table 3.4, the performance trends for the *MH-MPL* and *EM* methods are similar in the sense that *RBIAS* and *RRMSE* of the two methods get better as ω increases from 0.2 to 0.5, when estimating the parameter of the first copula θ_1 . For example, in the case of Frank-Gumbel mixture copula, the *RBIAS* associated to *MH-MPL* method decreases from 13 to 3.2% when ω increases from 0.2 to 0.5. This result is likely to be related to the fact that, when ω is small, the mixture model contains increasing information about one component but decreasing information about the other component. Indeed, a low mixing probability implies that the sampled data from the mixed model are close to the one component model. These results are in agreement with Fu *et al.* (2019)'s findings which showed that, in the

case of multivariate mixture distribution, an estimation method have a better performance when ω lies in the middle than on the boundaries of its space.

Table 3.4. RRMSE (%) and RBIAS (%) of different estimators, over 1000 replications, for different mixture models and different weights ω when sample size $n=150$ and dependence level $\tau = (0.6, 0.8)$

Copula		<i>EM method</i>				<i>MH-MPL method</i>			
		$\omega=0.2$		$\omega=0.5$		$\omega=0.2$		$\omega=0.5$	
		RRMSE	RBIAS	RRMSE	RBIAS	RRMSE	RBIAS	RRMSE	RBIAS
ω	C-C	35.7	19.8	17.4	-3.7	19.5	9.5	7.5	2
	F-F	38.9	12.4	14.7	-3.3	13.5	9.2	7.2	2.3
	G-G	34.8	-10	16.0	-4.2	16.5	8.1	7.9	1.8
	C-F	42.1	18.3	18.8	-3.8	17.1	8.2	7	1.4
	C-G	44.4	18.7	18.5	-4.8	14.1	9.4	7.4	1
	F-G	39.6	16	19.2	-4.7	17	8.3	7	1.8
Θ_1	C-C	59.4	-11.6	20.5	-4.2	29.1	14	9.4	2.8
	F-F	51.8	-15.9	20.7	7.3	29.8	15.2	8.5	3.1
	G-G	51.2	-14.4	22.2	5.7	28.4	13.4	8.8	3.2
	C-F	59.8	19.1	20.5	-5.4	27.1	11.7	8.3	2
	C-G	65.2	16.7	22.7	-6.4	28.7	10.9	7.8	1.9
	F-G	64.7	14.7	21.4	-4.8	27	13	8.3	3.2
Θ_2	C-C	10.2	2.1	20	-2.5	9.2	0.6	9.7	2.5
	F-F	11.2	1.6	19.5	4.2	9.4	1.5	9.8	2.4
	G-G	10.7	1.8	20.2	3.9	8.5	1.1	8.8	1.7
	C-F	10.1	1.9	22.8	-3.5	7.6	1.1	9.1	2.2
	C-G	9.9	2.1	23.4	3.7	7.3	0.9	7.9	2.3
	F-G	11.2	2.3	21.3	3.9	8.3	0.9	9.2	1.8

This table presents the RBIAS and RRMSE for different mixture of copula model and different weight ratio, using the metaheuristic maximum pseudo-likelihood (MH-MPL) and EM methods. Light gray color corresponds to best results regarding RBIAS values and dark gray color to best results regarding RRMSE.

However, it is also worth noting that the two methods exhibit noteworthy differences, especially under small ω . Indeed, in this case, the *MH-MPL* method leads to the best results with the minimum values of absolute *RBIAS* and *RRMSE*. For example, for the Clayton-Frank mixture model, the mean absolute *RBIAS* for the *MH-MPL* and *EM* methods are 19.1 and 11.7%, respectively, when

estimating θ_1 for $\omega=0.2$. These results imply that the *MH-MPL* estimates are closer to the true values of the parameters than those of the *EM* method. Accordingly, the *MH-MPL* method accurately estimates the parameters for low mixing probabilities ω and outperforms the *EM* method for both homogenous and heterogeneous mixture copulas. Moreover, according to Figure 3.4, results reveal that *MH-MPL* method yields more accurate estimates of the copulas' parameter θ_1 and θ_2 . Indeed, in all cases, *MH-MPL* method is characterised by smaller variability compared to *EM* method.

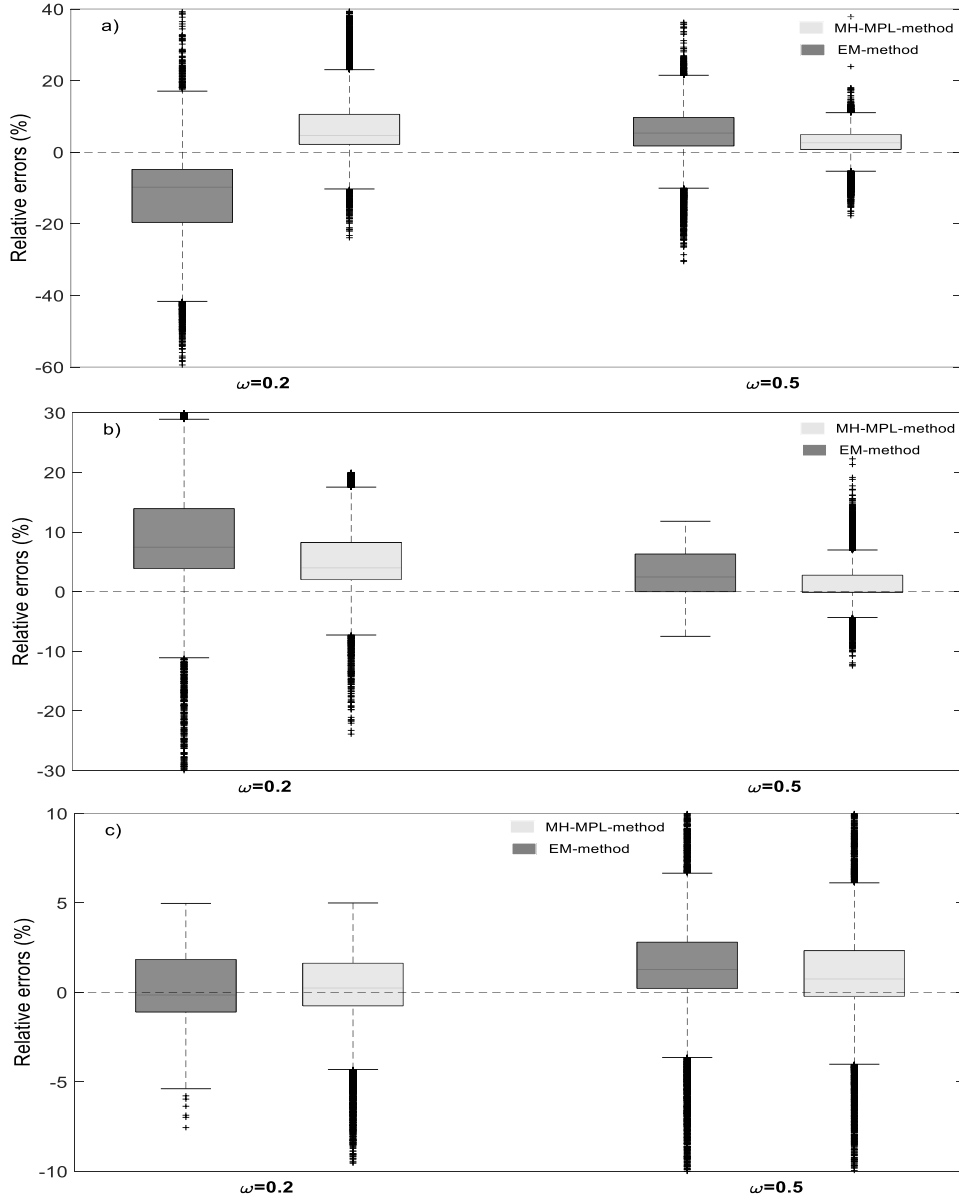


Figure 3.4. Relative errors results of MH-MPL and EM estimators for different weight ratio ω , dependence level $\tau = (0.6, 0.8)$ and $n=150$: a) weight ratio ω , b) first copula parameter θ_1 and c) second copula parameter θ_2

5. Conclusion

The main objective of the current study is to investigate the applicability of mixture copulas to model extreme events of hydrometeorological variables. Once the copula model and the associated marginal distributions are selected, parameter estimation is the key step. A novel contribution on mixture copula and in particular on flood frequency analysis has been provided by proposing a new

parameter estimation method. This proposed approach, denoted as *MH-MPL*, is based on the idea of estimating parameters by maximizing the log pseudo-likelihood function using a meta-heuristic algorithm. A simulation study is carried out in order to evaluate the proposed method and to compare the performances with the widely used *EM* method. The obtained results show that the proposed *MH-MPL* method is more appropriate than the *EM* method in terms of *RRMSE* and *RBIAS*, especially for small sample sizes. Moreover, the effectiveness of the *MH-MPL* method in dealing with weakly dependent data makes it particularly suitable for some hydrological applications.

In the current study, because of the associated advantages, the genetic algorithm is used for the optimization of the *MPL*. Future investigations will deal with other metaheuristic algorithms such the particle swarm optimization or by using a hybridization between two or several meta-heuristic algorithms. In addition, the extension of our approach to mixtures with more than two components is valuable and interesting for further research efforts.

**Chapitre 4. Multivariate L-moment based tests for copula selection,
with hydrometeorological applications**

Titre en français : Tests d'adéquation pour les copules, basés sur les L-moments multivariées, applications aux séries hydrométéorologiques

Auteurs :

Imene Ben Nasr Étudiante au doctorat à l'Institut National de la Recherche Scientifique (Centre Eau Terre Environnement), Université du Québec, 490 Rue de la Couronne, Québec, QC, Canada, G1K 9A9, Téléphone: (418) 261-6864, courriel: imene.ben_nasr@ete.inrs.ca

Fateh Chebana Professeur, Institut national de la recherche scientifique (Centre Eau Terre Environnement), Université du Québec, 490 Rue de la Couronne, Québec, QC, Canada, G1K 9A9, Téléphone: (418) 654-2542, courriel: Fateh.Chebana@ete.inrs.ca

A3. I. Ben Nasr and F. Chebana (2019). "Multivariate L-moment based tests for copula selection, with hydrometeorological applications." Journal of Hydrology 579:124151. DOI: <https://doi.org/10.1016/j.jhydrol.2019.124151>.

Dans cet article, I. Ben Nasr a développé deux nouveaux tests d'ajustement pour la sélection des copules multiparamètres dans un contexte hydrologique. I. Ben Nasr a présenté une comparaison entre les nouveaux tests et les tests existants. Tout au long de ce travail, F. Chebana a discuté l'aspect méthodologique et les résultats obtenus, et a révisé la version finale du manuscrit.

Abstract

Hydrological and climatological extreme events are characterized by several correlated random variables. For a better associated risk assessment, the dependence structure between these variables must be taken into account by considering copulas. Multiparameter copulas (M -copulas) play an important role by their flexibility and ability to capture more than one type of dependence. Since model misspecification can lead to underestimation or overestimation of the associated risk, it is of high importance to select the appropriate copula. To this end, several goodness-of-fit (GOF) tests have been proposed. However, these tests are applied, validated and evaluated for the usual one-parameter copulas. Nevertheless, there is no specific GOF test for M -copulas that takes into account the specificities of hydrometeorological series. Hence, the aim of the present paper is to introduce new GOF tests specifically for M -copulas and adapted to hydrometeorological context. More precisely, the proposed GOF tests are based on multivariate L-moments. A simulation study is conducted to evaluate and compare the performances of the proposed tests. The results confirm the usefulness of the new GOF tests in comparison with some well-established ones. Finally, the newly introduced tests are illustrated on hydrometeorological data series.

Keywords: Multivariate L-moments, flood, precipitation, parameter estimation, goodness-of-fit, power, multiparameter copula, testing

1. Introduction and literature review

Hydrometeorological events are characterized by a number of correlated random variables (e.g. Chebana and Ouarda, 2009; Hao and Singh, 2016). Therefore, a multivariate frequency analysis that takes into account the dependence between variables is of fundamental importance. Copulas have been proposed to model the dependence structure of these variables (e.g. Chebana, 2013; Genest and Chebana, 2017; Salvadori *et al.*, 2007). In hydrometeorology, copulas have gained substantial and increasing attention (e.g. Hao and AghaKouchak, 2013; Kao and Govindaraju, 2008; Salvadori and De Michele, 2004; Salvadori *et al.*, 2007; Vandenberghe *et al.*, 2010).

The usual one-parameter copulas have been largely developed in statistics and widely used in applications. However, multiparameter copula (denoted *M*-copula) is a developing topic with several open issues. Recently, Sadegh *et al.* (2017) et De Michele *et al.* (2013) used *M*-copulas to model meteorological drought events. In regional hydrological frequency analysis (HFA), Requena *et al.* (2016) used two-parameter Archimedean copulas to generate synthetic homogenous regions. With more flexibility, as shown by Salvadori and De Michele (2010), *M*-copulas are likely to better model dependence between multivariate data than classical one-parameter copulas.

In recent years, formal goodness-of-fit (*GOF*) tests have been introduced to select appropriate copula (e.g. Fermanian, 2005; Genest *et al.*, 2006; Mesfioui *et al.*, 2009 and references therein; Wang and Wells, 2000). These *GOF* tests could be *conceptually* valid for any copula structure. However, they are evaluated and validated only for one-parameter copula. Furthermore, as pointed out by Berg (2009), based on simulation results conducted for existing *GOF* tests, those based on empirical copula and Kendall process are recommended. Even for one-parameter copulas, these two recommended tests have some important drawbacks (as detailed in section 2.2). In addition, even though the importance and usefulness of general *GOF* tests, specific ones for given classes of

copulas are also required. This is in analogy to *GOF* tests for distributions in the univariate framework such as for normal distribution (e.g. Choulakian and Stephens, 2001; Chowdhury *et al.*, 1991; Zardasht *et al.*, 2015). Regarding copulas, Genest *et al.* (2011a) proposed a specific *GOF* test for Extreme-Value copulas whereas Durocher and Quessy (2017) focused on spatial copulas. When dealing with *M*-copula, one can find the test proposed by Kojadinovic and Yan (2011). The latter focuses only on Gaussian and Student copula as *M*-copulas. In addition, this test is exactly the empirical copula *GOF* test by Genest *et al.* (2006) where only the p-value evaluation is based on multiplier approach. They showed that the test is more appropriate for high dimensional and large sample size series, which is generally not the case of hydrometeorological series.

The main objective of the present paper is to introduce and evaluate new and specific *GOF* tests for *M*-copulas (see Table 4.1) based on multivariate L-moments. An extensive simulation study, involving a large number of *M*-copulas and dependence conditions, is conducted to evaluate the performance of the newly proposed tests and to compare with classical tests.

The paper is organized as follows. In Section 2, a short discussion of the theoretical background is presented (*M*-copulas, existing *GOF* tests and multivariate L-moments). The proposed multivariate L-moments based *GOF* tests are presented in Section 3 whereas Section 4 deals with the simulation study. Application to real-world hydrometeorological data is represented in Section 5. Concluding remarks are presented in the last section.

Table 4.1. Summary of goodness of fit tests, for hydrological applications

Copula	Available <i>GOF</i> tests	Drawbacks	References	Overall Comparisons
One-parameter	Kendall's Process	*Not suitable for small samples *Two different	Genest <i>et al.</i> (2006)	Genest <i>et al.</i> (2009) Berg (2009) Fermanian (2013)
	Moment based <i>GOF</i>	copulas might have the same Kendall 's function	Berg and Quessy (2009)	
	Empirical Copula Process	*Not appropriate for low dependence level	Fermanian (2005) Genest <i>et al.</i> (2006)	
Multiparameter	Empirical Copula Process based on multiplier approach for p-value estimation	* Focuses only on Gaussian and Student copula as M-copulas. *More appropriate for high dimensional and large sample size	Kojadinovic and Yan (2011)	The specific aim of the present paper
	Nonspecific test for general M-copula	-	The specific aim of the present paper	

2. M-copula in modeling hydrological variables

In this section a brief overview of *M*-copulas is presented (for more details see e.g. Joe (2014)). In the following, for the sake of brevity and simplicity, only the bivariate case is considered.

2.1. M-copulas

Extreme Value and Archimedean *M*-copulas are of special interest in hydrology and climatology (e.g. AghaKouchak, 2014; Salvadori and De Michele, 2010; Zhang and Singh, 2007b). Indeed, *M*-copulas can cover a range of dependence from perfect positive dependence to independence and may be extended to perfect negative dependence. *M*-copulas introduce additional flexibility to the model since they include other subfamilies of classical copulas as especial cases (e.g. De Michele *et al.*, 2013; Salvadori and De Michele, 2010). For instance, the *BB5* as 2-copula includes Gumbel

and Galambos ones. Consequently, M-copulas provide a better way to capture different mutual dependencies (e.g. Chen and Khashanah, 2014). Hence, this may improve the modelling features of the dependence structure.

A number of Archimedean M -copulas are available in the literature such as $BB1$, $BB6$ and $BB7$ (e.g. Joe, 2014). Besides, extreme-value copulas, which include $BB5$ as an example in the M -copula category, have been extensively used in recent years, given their importance in modeling catastrophic events (e.g. Salvadori and Michele, 2011; Zhang and Singh, 2012). For more details about properties of M-copulas including those used in this study (see Table 4.2), the reader can refer to Joe (2014).

Table 4.2. M-copulas used in this study (Joe, 2014).

Copulas	$C_{\vec{\theta}}(u, v)$	$\vec{\theta}$ space
BB1	$\left\{1 + \left[(u_1^{-\theta_1} - 1)^{\theta_2} + (v_1^{-\theta_1} - 1)^{\theta_2}\right]^{\frac{1}{\theta_2}}\right\}^{\theta_2}$	$\theta_1 \geq 0; \theta_2 \geq 1$
BB5	$\exp\left\{-\left[(-\ln u)^{\theta_1} + (-\ln v)^{\theta_1} - ((-\ln u)^{-\theta_1\theta_2} + (-\ln v)^{-\theta_1\theta_2})^{\frac{-1}{\theta_2}}\right]^{\frac{1}{\theta_1}}\right\}$	$\theta_1 \geq 1; \theta_2 > 0$
BB6	$1 - \left(1 - \exp\left(-\left[\{-\log(1 - u_1^{\theta_1})\}^{\theta_2} + \{-\log(1 - v_1^{\theta_1})\}^{\theta_2}\right]^{\frac{1}{\theta_2}}\right)^{\frac{1}{\theta_1}}\right)$	$\theta_1 \geq 1; \theta_2 \geq 1$
BB7	$1 - \left(1 - \left[(1 - u_1^{\theta_1})^{-\theta_2} + (1 - v_1^{\theta_1})^{-\theta_2} - 1\right]^{\frac{1}{\theta_2}}\right)^{\frac{1}{\theta_1}}$	$\theta_1 > 0; \theta_2 > 0$

2.2. Existing copula *GOF* tests

As indicated above, several *GOF* tests have been already proposed for copula selection. In the following, we focus on the tests described by Genest *et al.* (2009) as “Omnibus tests” and

recommended by Berg (2009), especially those based on the empirical copula process and the Kendall process. In fact, the other existing tests entail many shortcomings (Berg, 2009). For example, they involve many arbitrary choices (e.g. kernel type, window length, weight function) that make their application cumbersome. Furthermore, they are computationally expensive for high dimensional cases. The performance of the omnibus tests has not yet been evaluated when dealing with M-copulas. To the best of our knowledge, this study is the first one to present a critical review and to evaluate/compare their power in the case of M-copulas (Table 4.1).

The empirical copula based test proposed by Genest *et al.* (2009) consists in comparing the distance between the empirical copula C_n and an estimation C_{θ_n} of C obtained under the null hypothesis. Analogously to the test based on the empirical copula, Genest and Rivest (1993) proposed the Kendall-based test based on comparing the distance between the empirical Kendall's function K_n and an estimation K_{θ_n} of K obtained under the null hypothesis. In the following, we used the Cramér-von Mises based versions of these tests, denoted respectively by S_n^C and S_n^K .

When dealing with one-parameter copula, these two *GOF* tests have some drawbacks. For instance, their application is not suitable for small samples or low dependence levels (Berg and Quessy, 2009; Genest *et al.*, 2006). In addition, as pointed out by Genest *et al.* (2009), the Kendall-based test is not consistent for extreme-value (*EV*) copulas, since two different *EV* copulas have the same Kendall's function (Barbe *et al.*, 1996). Therefore, the test is not appropriate to distinguish between different *EV* copulas (e.g. Fermanian, 2013). In addition, as pointed out in Genest *et al.* (2009), the empirical copula-based *GOF* test is not appropriate for small sized samples, which is the case of the most of hydrological series (e.g. Laio *et al.*, 2009).

2.3. Multivariate L-moments

The multivariate L-moments were introduced by Serfling and Xiao (2007). They provide a summary and a description of the properties and shapes of a multivariate distribution. This makes them particularly useful in parameter estimation and hypothesis testing. Since multivariate (also univariate) L-moments are much less biased than classical moments, they are used as meaningful replacements of classical moments in a wide variety of applications, mainly in hydrology, climatology and meteorology analysis (e.g. Brahimi *et al.*, 2015; Chebana and Ouarda, 2007; Hosking and Wallis, 1993; Kysely and Picek, 2007).

The multivariate L-moments are defined as bellow

$$\lambda_{k[ij]} = cov(X^{(i)}, P_{k-1}^*(F_j(X_j))) \quad (4.1)$$

where COV is the covariance, k is the order of the multivariate L-moment, P_{k-1}^* is the shifted Legendre polynomial (Chang and Wang, 1983), $X^{(i)}$ is a random variable with marginal distribution F_i for $i = 1, 2$. Brahimi *et al.* (2015) introduced the copula's multivariate L-moments as

$$\lambda_k^{C[12]} = \int_0^1 \int_0^1 (C(u_1, u_2) - u_1 u_2) du_1 dP_k^*(u_2) \quad (4.2)$$

where C is the associated copula function and u_1, u_2 are the pseudo-observations. The sample version of the k^{th} copula multivariate L-moments is defined in term of pseudo-observations as (Brahimi *et al.*, 2015)

$$\hat{\lambda}_{k[12]}^C = \frac{1}{n} \sum_{i=1}^n \frac{R_i}{n+1} P_{k-1}^*\left(\frac{S_i}{n+1}\right) \quad (4.3)$$

where R_i is the rank of X_i^1 among $X_i^1, \dots, X_n^1$ and S_i is the rank of the concomitant $X_{i:n}^{[12]}$ among $X_1^2, \dots, X_n^2$. $X^{[12]}$ is formed by sorting X^2 in the ascending order and in turn shuffling X^1 by the

order of X^2 (Serfling and Xiao, 2007). The copula multivariate L-moments coefficient *ratios* are defined as

$$\tau_{k[12]}^{cop} = \frac{\lambda_{k[12]}^{cop}}{\lambda_{2[2]}}, \text{ for } k \geq 3 \text{ and } \tau_{2[12]}^{cop} = \frac{\lambda_{2[12]}^{cop}}{\lambda_{1[2]}} \quad (4.4)$$

According to Brahim *et al.* (2015), the vector of the multivariate copula L-moments converges in distribution to the multivariate normal distribution with variance-covariance matrix $\Sigma = (\sigma_{ij})_{1 \leq i, j \leq 4}$ defined as

$$\sigma_{ij} = b_i b_j - \frac{1}{n^{(i+j+2)}} \sum_{1 \leq k < l \leq n} \left[((k-1)^i (l-i-2)^j + (k-1)^j (l-j-2)^i) X_{k:n}^{(12)} X_{l:n}^{(12)} \right] \quad (4.5)$$

where $b_k = \left(\frac{1}{n^{k+1}} \sum_{i=1}^n (i-1)^k X_{[i:n]}^{(12)} \right)_{k=0, \dots, 3}$.

3. The proposed *GOF* tests for *M*-copulas

In this section, we present the new proposed *GOF* tests for *M*-copulas based on multivariate copula L-moments. Two nonparametric approaches of tests are developed (see Table 4.3).

Table 4.3. Summary of the proposed *GOF* tests for *M*-copulas

	Special cases from literature	Advantages
Multivariate L-moments based <i>GOF</i> Z_n^X	Hosking and Wallis (1993) <i>GOF</i> test for univariate distribution	-Consistent even for small sample -Outperform already existing tests -Characterize a wide range of copula -Does not depend on any
Parameters estimation based <i>GOF</i> Z_n^{LM}	Shih (1998) <i>GOF</i> test for Clayton copula Berg and Quessy (2009) moment-based <i>GOF</i> test	assumptions concerning the parent copula -Involves no choice for external factors

The first one is based on the distance between the multivariate L-moments of the fitted copulas and the sample multivariate L-moments. It is inspired by the Hosking and Wallis (1993) *GOF* test

proposed in the univariate setting. The second approach consists in generalizing the *GOF* test developed by Shih (1998). It is based on the deviation between two estimators of the parameters vector of the *M*-copula using the multivariate L-moments estimation method (Brahimi *et al.*, 2015) and Maximum Pseudo-Likelihood (*MPL*) estimation method (Kojadinovic and Yan, 2010).

3.1. Proposed *GOF* test based on multivariate copula L-moments

In this section, a *GOF* test for *M*-copulas, based on multivariate copula L-moments ratios, is presented. The main idea of the new developed *GOF* test consists in evaluating a “distance” between the empirical values of $\hat{\tau}_{3[12]}$ and $\hat{\tau}_{4[12]}$ and their theoretical counterparts (equation 4.4) associated to the *M*-copula C_{θ_n} under the null hypothesis H_0 . Let $\hat{\omega} = (\hat{\tau}_{3[12]}, \hat{\tau}_{4[12]})$ be the vector of the empirical multivariate L-moments ratios and $\omega^{cop} = (\tau_{3[12]}^{cop}, \tau_{4[12]}^{cop})$ be their theoretical counterparts from the underlying *M*-copula. The associated statistic is defined as

$$Z_n^X = n(\hat{\omega} - \omega^{cop})^t \Omega^{-1} (\hat{\omega} - \omega^{cop}) \quad (4.6)$$

where $\Omega = \left(\frac{\sigma_{ij}}{n\lambda_2^2} \right)_{(i,j=3,4)}$ and σ_{ij} are defined by equation (4.5).

It is worth mentioning that for traditional *GOF* tests, the asymptotic distributions of their statistics depend on the unknown copula $C_{\bar{\theta}}$ (Genest *et al.*, 2009) and their *p-values* could only be obtained via bootstrap procedures (Genest and Rémillard, 2008). Hence, for convenience and comparison purpose, the bootstrap procedure is used to estimate *p-values* of the proposed test. Asymptotic results are out of the scope of the present paper and could be the object of future work.

The use of the proposed statistic in equation (4.6) has several advantages, including:

- a. Simple formula available for the proposed statistic in terms of the ranks of the observations;

- b. The statistic involves no subjective choices, such as the choice of a kernel and associated smoothing parameters;
- c. Compared to conventional moments, multivariate L-moments are less subject to bias in estimation;
- d. Multivariate L-moments are able to characterize a wide range of M -copulas.

3.2. Extension of the Shih's *GOF* test to M -copulas

When dealing with the one-parameter Clayton copula, Shih (1998) proposed a “moment-based” *GOF* test. The basic idea of the test consists in the comparison between two estimators of the scalar parameter θ via Kendall's τ and a weighted rank-based estimator. Berg and Quessy (2009) extended this test for any one-parameter copula based on classical copula moments (Kendall's τ and Spearman's ρ). However, the case of M -copula has not been investigated. In this section, we propose an extension of the moment-based *GOF* test to deal with M -copulas. The proposed *GOF* test is, then, based on the comparison of two parameters vector estimates, namely, the multivariate L-moments estimator $\vec{\theta}_{Lmom}$ and the maximum pseudo-likelihood estimator $\vec{\theta}_{MPL}$. As in the tests by Shih (1998) and Berg and Quessy (2009), under the null hypothesis that the unknown copula of a population belongs to a specific family, both estimators converge to the true parameters. However, if the assumed copula is invalid, the two estimators generally do not converge to the same value (White, 1981).

In their simplest form, assume a vector of p parameters $\vec{\theta} = (\theta_1, \dots, \theta_p)^t$. Let $l_{[12]} = (\lambda_{1[12]}, \lambda_{2[12]}, \dots, \lambda_{p[12]})^t$ be the vector of the first p -variate L-moments of the fitted copula $C_n(\vec{\theta})$.

The estimator $\hat{\vec{\theta}}_{Lmom} = A^{-1}l_{[12]}^t$ is a consistent estimator of $\vec{\theta}$, where A is the coefficient of the shifted Legendre polynomial P_{k-1}^* (Brahimi *et al.*, 2015). We define the statistic Z_n^{LM} as

$$Z_n^{LM} = n \sum_{i=1}^p (\hat{\theta}_{Lmom,i} - \hat{\theta}_{MPL,i})^2 \quad (4.7)$$

Note that the statistic Z_n^{LM} is defined for any d -variate copulas. Indeed, as pointed out in Joe (2014), for a d -variate copula, the number of parameters is of the order of the dimension d up to the square of the dimension. As above, the corresponding p -value is estimated by bootstrap. The proposed *GOF* test, based on parameter estimation methods, has many advantages:

- a. Its statistic is free of any subjective choices, which could affect the performance of the tests;
- b. Compared to conventional moments, multivariate L-moments are less subject to bias in estimation;
- c. Parameter estimates obtained from multivariate L-moments are usually more accurate in small samples than the estimates obtained by other methods.

4. Simulation study

4.1. Simulation design

The objective of the simulation study is to evaluate the performance of the proposed *GOF* tests as well as the already existing ones when the latter are able to test M -copulas. Typically, the performance of a *GOF* test is evaluated through the estimation of first type error α (level of the test) as well as the test power (the rejection rates of the null hypothesis). In this study, α is fixed to 5%, as a usual value in several previous studies (e.g. Berg, 2009; Genest *et al.*, 2009).

It was pointed out in the literature (e.g. Berg, 2009; Fermanian, 2013; Genest *et al.*, 2006), in the case of one parameter copulas, that the performance of a *GOF* test is influenced by several factors, especially, sample size and dependence level. Hence, a sensitivity analysis is performed to identify the effect of these factors on the behavior of each *GOF* test. Since for most flood events Kendall's τ is between 0.3 and 0.8 (e.g. Requena *et al.*, 2013a; Zhang and Singh, 2007b) and for comparison

purpose, we considered $\tau=0.2, 0.4, 0.6$, as in previous *GOF* testing studies (e.g. Berg and Quessy, 2009).

Regarding the effect of sample size, $n = 30, 50, 100$ are investigated. The values of n are selected on the basis of situations commonly encountered in flood frequency analysis. Indeed, Barth *et al.* (2017) reported that data examined from 1375 long-term U.S. Geological Survey (USGS) stream gage sites have at least 30 years annual peak flow. Furthermore, Santhosh and Srinivas (2013) indicated that the minimum sample size drawn from five typical peak flow records, from different parts of the world, is equal to 40. For all dependence scenarios, $M = 10,000$ samples are generated. The latter is typically used in *GOF* testing simulations (e.g. Genest *et al.*, 2009).

Data are generated through Monte-Carlo simulations from representative and most used copulas in hydrometeorology analyses (e.g. Sadegh *et al.*, 2017; Salvadori *et al.*, 2007). Therefore, two classes of copula family are used: 1) Archimedean *M*-copulas including *BB1*, *BB6* and *BB7* and 2) Extreme-Value *BB5* copula. To generate data from a given *M*-copula, we applied the procedure proposed by Nelsen (2006), based on the *conditional distribution method*.

4.2. Simulation results

In this section, simulation results obtained by the application of different *GOF* tests to the considered copulas are presented. Comparisons are given in the last part of this section.

4.2.1. First type error estimates

The first desirable property of a *GOF* test is the ability to maintain the first type error close to the significance level. Probabilities of rejection are estimated under H_0 and compared to the theoretical nominal level. The effect of sample size on the level of considered *GOF* tests is shown in Table 4.4 for dependence level $\tau=0.6$.

Table 4.4. Level of the proposed *GOF* tests (in %) for various sample sizes and dependence level $\tau=0.6$

Copula under H_0	n=30		n=50		n=100	
	Z_n^χ	Z_n^{LM}	Z_n^χ	Z_n^{LM}	Z_n^χ	Z_n^{LM}
BB1	7.2	7.0	5.0	5.4	5.3	4.9
BB5	7.0	6.0	4.8	5.2	4.8	4.0
BB6	6.4	5.7	5.1	5.0	4.2	4.7
BB7	5.9	6.3	5.4	4.9	5.0	5.1

This table presents the obtained first type errors, for various sample sizes, for each test and several kinds of copulas. Values in bold indicate the optimal first type error.

For $n=50$ and $n=100$, it is found that the new developed *GOF* tests perform well in terms of first type error since its estimation is very close to $\alpha = 5\%$ (between 4 and 5.2%). For sample size $n=30$, the two proposed tests Z_n^{LM} and Z_n^χ have slightly large first type errors (between 5.7 and 7.1%). This is expected since the sample L-coskewness and L-cokurtosis ratios, $\hat{\tau}_{3[12]}$ and $\hat{\tau}_{4[12]}$ respectively, can be slightly biased for small n (Serfling and Xiao, 2007). In addition, this might be explained also by a tendency for $\hat{\sigma}_{33}$ and $\hat{\sigma}_{44}$ (estimated by equation (4.5)) to underestimate the true variance of L-coskewness and L-cokurtosis, for small n , as in the univariate case (Kjeldsen and Prosdocimi, 2015b).

Table 4.5 provides results of application of the proposed tests for sample size $n=100$ and different dependence levels.

Table 4.5. Level of the proposed *GOF* tests (in %) for various dependence levels and sample size $n=100$

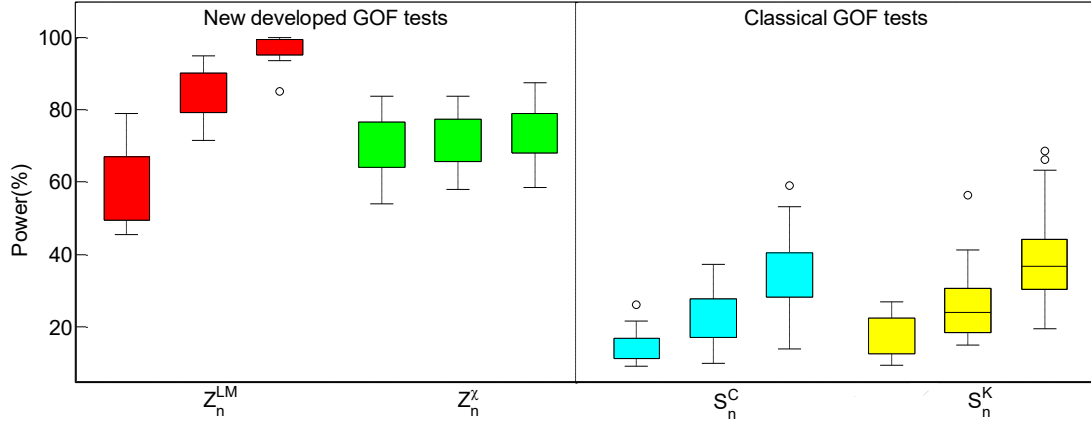
Copula under H_0	$\tau=0.2$		$\tau=0.4$		$\tau=0.6$	
	Z_n^χ	Z_n^{LM}	Z_n^χ	Z_n^{LM}	Z_n^χ	Z_n^{LM}
BB1	10.0	7.8	6.0	5.3	5.3	4.9
BB5	8.3	6.2	6.2	5.0	4.8	4.0
BB6	9.1	6.0	7.1	5.0	4.2	4.7
BB7	7.0	7.2	7.0	5.1	5.0	5.1

This table presents the obtained first type errors, for various dependence levels, for each test and several kinds of copulas. Values in bold indicate the optimal first type error.

From Table 4.5, one can see that the first type errors of the proposed tests Z_n^{LM} and Z_n^χ appear in general to be closer to 5% as the dependence level increases which is in agreement with other studies (e.g. Durocher and Quessy, 2017; Genest *et al.*, 2009). Indeed, for low dependence level $\tau=0.2$, first type errors are between 6 and 10%. This is expected since for low dependence level, the distinctive features of the models are fuzzier (e.g. Genest *et al.*, 2009). As a consequence, different tests fail to distinguish between different copulas (e.g. Berg, 2009; Berg and Quessy, 2009; Mesfioui *et al.*, 2009). However, the higher dependence level is ($\tau=0.6$), the better the first type errors are (between 4 and 5.3%). For moderate dependence level ($\tau=0.4$), the Z_n^χ test fails to hold the nominal level (estimated levels above 6%). This is likely due to the fact that the copula L-moments are approximatively equal to zero for near-independently data (Brahimi *et al.*, 2015; Serfling and Xiao, 2007).

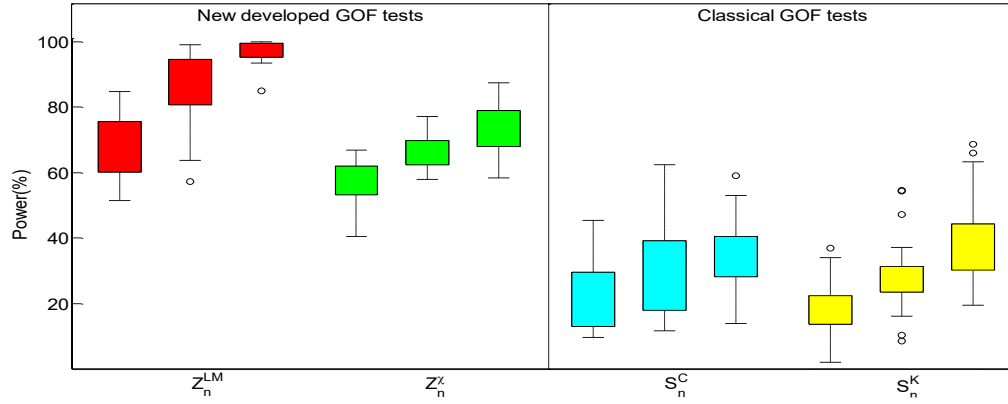
4.2.2. Power evaluation

In order to be more informative, the powers of the considered tests are presented as boxplots in Figures 4.1 and 4.2.



Each color corresponds to a test. The first box corresponds to $n=30$, the second for $n=50$ and the third for $n=100$.

Figure 4.1. Boxplots showing the power of different *GOF* tests various sample sizes and dependence level $\tau=0.6$



Each color corresponds to a test. The first box corresponds to $\tau=0.2$, the second for $\tau=0.4$ and the third for $\tau=0.6$.

Figure 4.2. Boxplots showing the power of different *GOF* tests various dependence levels and sample size $n=100$

From Figure 4.1, one can see that, in most cases, the proposed test Z_n^{LM} shows almost the highest power for several M -copulas. The power of the Z_n^{LM} is between 85 and 100% for sample size $n=100$ and dependence level $\tau=0.6$. Even for a moderate sample size ($n=50$), Z_n^{LM} has the highest power. In the same way, the power of the test statistic Z_n^X increases with the sample size. Indeed, large sample size not only helps to distinguish between copulas but also plays a role in the reliability of the bootstrap procedures used to approximate the associated p -values as it is the case for one parameter copula (e.g. Berg, 2009; Berg and Quessy, 2009). Moreover, as pointed out in the

literature of *GOF* testing, when sample size is small, it is more difficult to distinguish between copulas.

Since the dependence level affects the power of a test in a similar way as the sample size, the following discussion will be brief. From Figure 4.2, for low dependence level $\tau=0.2$, it is found that results of the developed *GOF* tests are satisfactory. Indeed, Z_n^{LM} has the highest power (between 51 and 85%) and Z_n^X power estimates are between 40 and 69%. From Figure 4.2, it is found that the power of the proposed *GOF* tests appear to improve as the dependence level increases. The same pattern is observed in *GOF* testing literature. This can be explained by the fact that, when dependence level increases, the copula structure is more distinguishable (e.g. Berg, 2009). However, when the dependence level is too small, the dependence structures are hardly distinguishable.

4.2.3. Comparisons

A comparison of the performance between the proposed tests and the classical ones is made through boxplots (Figure 4.1 and 4.2). According to Figure 4.1, the power estimates of all statistics increase when n becomes larger. This agrees with similar findings by Berg (2009) and Genest *et al.* (2009) made for one-parameter copula *GOF* tests. However, Z_n^X is less influenced by n with only a difference of about 2% when n increases from 30 to 100. This can be explained by the fact that multivariate L-moments are robust and suffer less from the effects of sampling variability (Brahimi *et al.*, 2015). Nevertheless, Figures 4.1 and 4.2 show that the two classical statistics S_n^C and S_n^K are not able to distinguish the different *M*-copulas. The associated powers of these tests are between 9 and 27% for $n=30$, and between 20 and 69%, for $n=100$. For S_n^K this is unsurprising since the Kendall function could be the same for two different copulas (e.g. Nelsen, 2006).

On the other hand, S_n^C behaves in the same way as in the case of one-parameter copula. For M -copula, low power estimates may be caused by the small sample size ($n < 150$) (e.g. Genest *et al.*, 2006). The most striking result, from Figure 4.2, is that the tests S_n^C and S_n^K have lower powers even for strong dependence ($\tau = 0.6$), when dealing with M -copulas. A possible explanation, in the case of the statistic S_n^K , might be that as the number of parameters increases, the univariate summary represented by the probability integral transformation and its distribution function K is less representative of the multivariate dependence structure. In the case of S_n^C , it might be related to the fact that the sample size is relatively small (even $n = 100$), and as justified by Genest *et al.* (2006), this test has high power for $n > 150$.

From Figure 4.1, the different *GOF* tests can be ranked in term of power as $Z_n^{LM} > Z_n^X > S_n^K > S_n^C$. Overall, for a given n and τ , the proposed tests have higher powers than those of the classical tests. In fact, since the shape of copula structure can be described by their multivariate L-moment matrices (Serfling and Xiao, 2007), the proposed statistics allow a better differentiation of various copulas. Note that, the *GOF* test Z_n^{LM} based on parameters estimation method has always the best performance. This may be due to inconsistency of the test. In fact the parameter estimates by the *MPL* and multivariate L-moments may be identical, as supported by Berg and Quessy (2009) in the case of moments-based *GOF* test, and so the test will accept more often the null hypothesis.

One wants to study the performance of the proposed *GOF* tests when dealing with one-parameter copulas. Table 4.6 provides corresponding results obtained from the proposed tests and the classical ones. The proposed tests Z_n^{LM} and Z_n^X perform well in terms of power compared to the classical tests. The power estimates for the proposed tests are between 45.0 and 87.0%. For classical *GOF* tests power estimates are between 33.4 and 84.3%. Noticeable exception is made in the case of Z_n^X statistic, which is outperformed by the classical tests, when dealing with the first type errors.

Table 4.6. Performances of proposed and classical *GOF* tests (in%) in the case of one parameter copulas for sample size $n=100$ and dependence level $\tau=0.6$

Copula under H_0	True copula	Proposed <i>GOF</i> tests		Classical <i>GOF</i> tests	
		Z_n^X	Z_n^{LM}	S_n^C	S_n^K
Clayton	Frank	65.3	75.2	70.5	73.0
	Clayton	8.2	4.8	7.2	5.8
	Gumbel	45.0	60.0	44.4	84.3
Frank	Clayton	65.9	87.0	56.3	62.4
	Frank	6.1	4.6	9.2	6.3
	Gumbel	67.0	88.5	35.0	33.4
Gumbel	Clayton	62.0	60.0	34.2	69.8
	Gumbel	6.2	5.3	5.0	4.8
	Frank	54.1	72.7	54.3	61.8

This table presents the obtained rejection rates at significance level $\alpha = 5\%$ for each test and several kinds of one parameter copulas. Rejection rates are expressed in percentages. Values in bold indicate the higher rejection rate and the empirical levels are italicized.

5. Illustrative case studies

After evaluating and comparing the proposed *GOF* tests on simulated data, in this section, we apply these tests on real-world hydrometeorological data. A comparison is also made with the classical Kendall and empirical *GOF* tests. The purpose of these applications is to illustrate the practical use of the proposed tests in order to select an appropriate copula family. The first application deals with flood events and the second with meteorological drought. Note that these two applications are characterized by different record lengths, dependence level and dependence shape.

5.1. Bivariate Flood Frequency Analysis

To scrutinize the impact of length of data on the performance of the *GOF* tests, we use two flood peak [Q (m^3/s)] and volume [V (m^3)] data series. The first one, from the Romaine station, located in the province of Quebec, Canada, for the period 1984-2010. The second data set corresponds to the Ankang hydrological station at the Upper Hanjiang River, China. The same series is used in Xiong *et al.* (2015), for the period 1950-2011. Since the focus of the study is on the dependence structure, the choice of marginal distributions is performed but not presented.

First, the dependence analysis between variables is carried out by dependence measures. In order to discard inappropriate copulas, the Kendall's τ is adopted to measure a quantitative value of the dependence structure. Kendall's τ is 0.30 and 0.8 respectively for Romaine and Ankang stations. For Ankang station, this value supports the positive and strong dependence between Q and V . However, for Romaine station, the dependence is relatively low. Hence, copulas with dependence range not matching with those of the series are discarded and four copulas are retained for subsequent analysis: M -copulas ($BB1$ and $BB5$) and one-parameter copulas (Gumbel and Galambos). The parameter of the considered copula is estimated using the Maximum-Pseudo likelihood method (e.g. Joe, 2014).

Besides quantitative measures of the dependence, graphical analysis of the dependence is conducted. Scatter plots and Kendall's function plots are used to evaluate if a chosen copula fits properly the observed data. A set of 10 000 synthetic pairs are generated from each one of the considered copulas. Then, the synthetic pairs are transformed by using univariate marginal distributions. A number of common univariate distributions in *HFA* are also considered to model each of the univariate series (including Weibull, Gumbel, normal, *GEV* and Gamma distributions). The Anderson-Darling *GOF* test (Laio, 2004) is used to select the accepted distributions for each variable. Then, the Akaike Information Criterion (*AIC*) is used to select the best distribution among the accepted ones for each variable (Akaike, 1974; Sadegh *et al.*, 2017). The selected distribution is the one associated to the lowest values of *AIC*. The corresponding results are presented in Table 4.7 for flood peak and volume series. These results are in agreement with those obtained in other studies (e.g. Chebana *et al.*, 2016; Rao and Hamed, 2000).

Table 4.7. Flood peak and volume results of the goodness-of-fit test and model selection criteria, for univariate distribution selection

Variable		Q			V			
Station	Accepted Distributions	Estimated Parameters	pvalue	AIC	Accepted Distributions	Estimated Parameters	pvalue	AIC
Romaine	Weibull	$\beta=3.85$ $\eta=1584$	0.59	313	GEV	$\beta=0.33$ $\eta=928$ $\theta=3205$	0.34	339
	Gumbel	$\mu=1294$ $\sigma=280$	0.98	307	Gumbel	$\mu=3233$ $\sigma=184$	10^{-6}	347
	Gamma	$\beta=18$ $\eta=81$	0.81	308	Gamma	$\beta=21.41$ $\eta=169$	10^{-16}	342
	Normal	$\mu=1454$ $\sigma=362$	0.61	311	Normal	$\mu=3617$ $\sigma=725$	0.51	340
	Weibull	$\beta=1.99$ $\eta=9950$	0.83	1219	Weibull	$\beta=2.03$ $\eta=1.8 \cdot 10^9$	0.94	2723
Ankang	Gumbel	$\mu=6432$ $\sigma=3986$	0.11	1250	Gumbel	$\mu=12.1 \cdot 10^8$ $\sigma=7.3 \cdot 10^8$	0.22	2751
	GEV	$\beta=0.20$ $\eta=4448$ $\theta=6837$	0.26	1234	GEV	$\eta=8.1 \cdot 10^8$ $\theta=1.2 \cdot 10^9$	1	2736
	Normal	$\mu=8645$ $\sigma=4532$	0.59	1223	Normal	$\mu=16 \cdot 10^8$ $\sigma=8.5 \cdot 10^8$	0.79	2729

β : shape parameter, η : scale parameter, θ : location parameter, μ : mean and σ : standard deviation. The parameters of each distribution are estimated via Maximum likelihood method except for the GEV distribution, the parameters are estimated via L-moments method. Bold value corresponds to the best fitting distribution.

From Figure 4.3a, presenting scatter plots for Romaine station, we can see that Extreme-Value copulas (i.e. *BB5* and Galambos) are sharper in the upper right corner than Archimedean copulas (*BB1* and Gumbel) which are more scattered. This figure shows also the positive lower tail dependence of the *BB1* copula. We can conclude that Extreme-Value copulas discard the largest observations whereas the *BB1* copula covers these observations in the generated sample, as its dependence structure is more spread in the upper tail.

Figure 4.4 presents K-plots for different copulas. Since extreme-value (*EV*) copulas entail the same Kendall's function, the K-plots cannot distinguish among Galambos and *BB5* copulas and provide the same information for these two *EV* copulas. Note that the Kendall's curves are identical for the Galambos and *BB5* copulas (dash-dotted line). As expected, we can see from Figure 4.4a that the distance between the empirical and theoretical Kendall's function is greater for Extreme-Value copula than for Archimedean copula. This analysis shows that Extreme-Value copulas are not

suitable for fitting Romaine data. Consequently, *BB1* copula seems to be the best fitting copula. However, a further analysis is needed to confirm this preliminary conclusion.

To select the best fitting copula in a formal way, the above described *GOF* tests are applied. *GOF* test statistics and associated *p-values* are summarized in Table 4.8. It can be seen that *BB1* copula is accepted by the proposed tests Z_n^{LM} and Z_n^X but rejected by the classical tests S_n^C and S_n^K . The Gumbel copula is accepted by proposed tests as well as S_n^K . However, S_n^C rejects all copulas except the Galambos one. Indeed, as indicated above, this test is affected by small sample size ($n=26$).

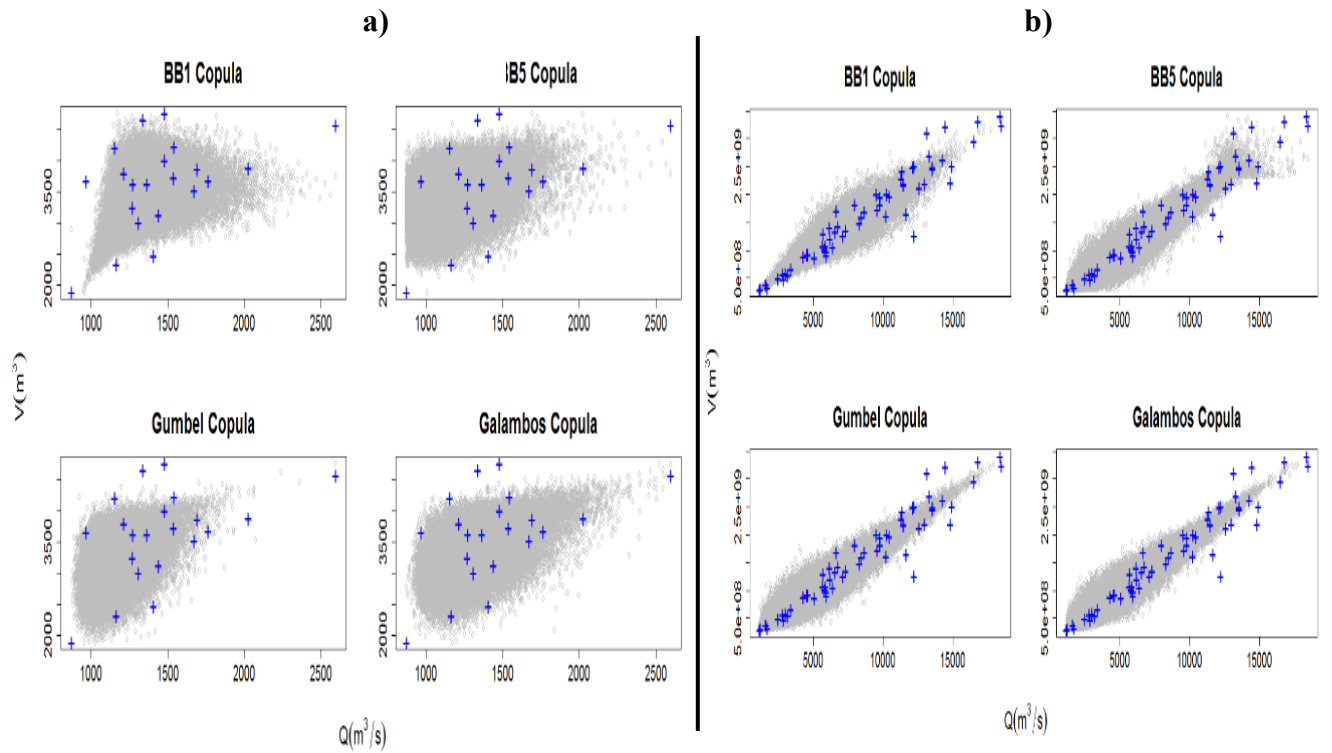


Figure 4.3. Scatter plot of observed data and 10 000 generated pairs from the copula fitted: a) for the Romaine station and b) for the Ankang station

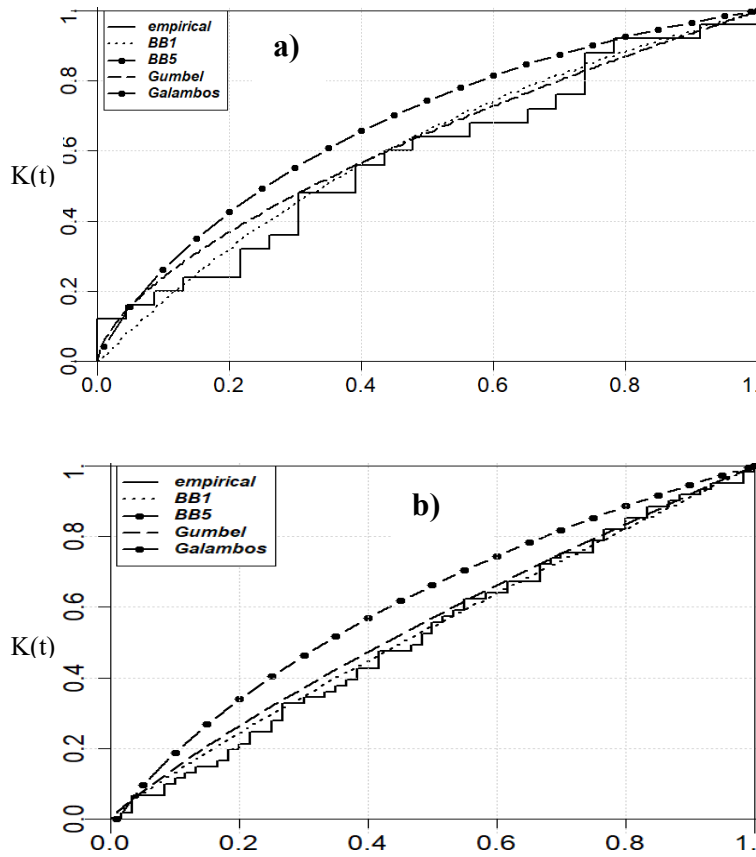


Figure 4.4. Comparison between the empirical estimate and theoretical Kendall's function for each fitted copula for (a) Romaine data and (b) Ankang data

Table 4.8. Results of the application of the *GOF* tests: Estimated copula parameters by Maximum Pseudo-Likelihood method and corresponding *p-value* based on 1000 parametric bootstrap samples, for a significance level $\alpha = 5\%$, for Romaine and Ankang data series

Stations	Copula family	Copula parameters		Proposed <i>GOF</i> tests (p-values)		Classical tests (p-values)		AIC
		θ_1	θ_2	Z_n^X	Z_n^{LM}	S_n^C	S_n^K	
Romaine	BB1	1.09	1.00	0.42	0.58	0.04	0.03	-28.6
	BB5	0.27	1.17	0.01	0.07	0.02	0.05	-26.0
	Gumbel	1.43	-	0.63	0.44	0.03	0.05	-24.7
	Galambos	0.73	-	0.03	0.05	0.06	0.05	-11.7
Ankang	BB1	3.00	2.44	0.06	0.09	0.01	0.03	-168.8
	BB5	212.56	0.02	0.62	0.59	0.53	0.05	-136.3
	Gumbel	4.97	-	0.04	0.06	0.03	0.05	-138.3
	Galambos	4.24	-	0.50	0.63	0.56	0.05	-137.9

This table presents the estimated p-values at the significance level $\alpha = 5\%$ for each test. Values in bold indicate the copulas that pass the *GOF* tests. The grey color indicates the best fitting copula regarding the model selection criterion (AIC).

It is apparent from Table 4.8 that S_n^K is unable to distinguish between different Extreme-Value copulas (Gumbel, *BB5* and Galambos). Indeed, regarding S_n^K , all considered Extreme-Value copulas are accepted as a good candidate for fitting the Romaine data, as they have the same Kendall's function (Genest and Boies, 2003). For Z_n^{LM} , all copulas are also accepted. This can be explained by the inconsistency of the test. Therefore, it is important to bear in mind that the parameter estimates by the *MPL* and multivariate L-moments may be identical even if the copula does not fit well the data (e.g. Berg and Quessy, 2009). In order to select the best fitting copula among those accepted by different *GOF* tests, the model selection criterion *AIC* is considered (Sadegh *et al.*, 2017). Accordingly, the *BB1* copula is selected as the most appropriate copula for (Q, V) of the Romaine station. As a reminder, the *BB1* copula is accepted by the new proposed tests whereas it is rejected by the classical tests.

In the case of the Ankang data, as shown in Figure 4.3b, one-parameter copulas (Gumbel and Galambos copula) and the two-parameter *BB5* copula (contains the Gumbel and Galambos copulas) display rather similar behaviour in terms of scatter plots. These copulas reproduce suitably the dependence in the upper extremes. Moreover, *BB1* copula reproduces suitably the dependence in the upper and lower tails since its structure is more spread. From Figure 4.4b, theoretical Kendall's function of the *BB1* and Gumbel copulas are closer to the empirical Kendall's function than for Extreme-Value copulas. Then, *BB1* and Gumbel copulas could better fit the Ankang data.

As shown in Table 4.8, *BB1* copula is accepted by the proposed tests Z_n^{LM} and Z_n^χ while it is rejected by the classical tests S_n^K and S_n^C . *BB5* and Galambos copulas are accepted by all *GOF* tests. However, Gumbel copula is accepted by Z_n^{LM} and S_n^K while it is rejected by Z_n^χ and S_n^C . Note that, for Ankang station, compared to Romaine station, the sample size is large and dependence level is

high. Hence, it would be easier to distinguish between the different copulas. According to *AIC*, the *BB1* copula could be selected as the best one to fit data.

These applications show, in a practical and complimentary way to the simulation results, the advantage of using the new proposed tests especially when dealing with small sample size and low dependence level. In addition, these applications confirm the ability of the proposed tests to distinguish between different *EV* copulas with respect to a Kendall approach, which necessarily shrinks into the parameter all the information about the *EV*-copula.

The main objective of this study is to deal with hydrological series characterized by small sample size. As pointed out through the simulation study and the application to real-world data, this objective is reached with the proposed tests.

5.2. Bivariate Drought analysis

In this section, we focus our attention on drought analysis with an emphasis on the use of *GOF* tests for the best fitting copula selection. Since drought is one of the most damaging natural hazards, its prediction is of high importance for risk assessment and decision making. Droughts can be of meteorological, agricultural, hydrological, or socio-economical nature. In this study, we focus on meteorological droughts caused by a deficit in precipitation and a lack of soil moisture. Hence, we employ the multivariate standardized drought index (*MSDI*) proposed by Hao and AghaKouchak (2013) where a sequence of negative *MSDIs* indicates a drought event.

In this application, we used monthly precipitation and soil moisture data from 1948 to 2017, obtained from the Climate Prediction Center (*CPC*), for climatic division # 1 in northern California. Different copulas are used to construct the joint probability distribution of precipitation (mm/day) and soil moisture (*mm*) anomalies (defined as deviations from the 1971-2000 monthly

climatology). To scrutinize the impact of length of data on the dependence structure and the performance of the underlying *GOF* tests, we use a subset of recent 30 years (1988–2017) of annual data, in addition to the original 70 years (1948–2017) of record. Similar procedure, as for flood variables, is also applied to each of the variable precipitation and soil anomalies series. The obtained results are summarized in Table 4.9. These results are consistent with those of other studies (e.g. Khedun *et al.*, 2014).

Table 4.9. Precipitation and soil moisture anomalies results of the goodness-of-fit test and model selection criteria, for univariate distribution selection

Variable	Precipitation anomalies				Soil moisture anomalies			
Sample size	Accepted Distributions	Estimated parameters	pvalue	AIC	Accepted Distributions	Estimated parameters	pvalue	AIC
n=30	Gumbel	$\mu=-109.6$ $\sigma=232$	0.26	416	GEV	$\beta=-0.65$ $\eta=564$ $\theta=41$	0.29	456
	GEV	$\beta=0.49$ $\eta=217$ $\theta=-52.6$	10^{-7}	10^{10}	Gumbel	$\mu=-209$ $\sigma=740$	0	485
	Normal	$\mu=-2.67$ $\sigma=201$	10^{-5}	407	Normal	$\mu=101.2$ $\sigma=544.3$	10^{-7}	467
n=70	Gumbel	$\mu=-142$ $\sigma=254$	0.12	988	Gamma	$\beta=226$ $\eta=0.22$	0.91	10.4
	GEV	$\beta=0.19$ $\eta=247$ $\theta=-114$	0.19	1023	Gumbel	$\mu=-323$ $\sigma=675$ $\beta=0.59$	0	10^{11}
	Normal	$\mu=-13.24$ $\sigma=265.8$	0.72	984	GEV	$\eta=677.8$ $\theta=-122.6$	$4 \cdot 10^{-3}$	10^{10}
					Normal	$\mu=-0.86$ $\sigma=611$	$2 \cdot 10^{-5}$	1100

β : shape parameter, η : scale parameter, θ : location parameter, μ : mean and σ : standard deviation. The parameters of each distribution are estimated via Maximum likelihood method except for the GEV distribution, the parameters are estimated via L-moments method. Bold value corresponds to the best fitting distribution.

Figure 4.5 shows the scatter plots and different fitted copula to describe the dependence structure between precipitation and soil moisture anomalies using 70 years of annual (Figure 4.5a), and 30 years of annual (Figure 4.5b) observed data.

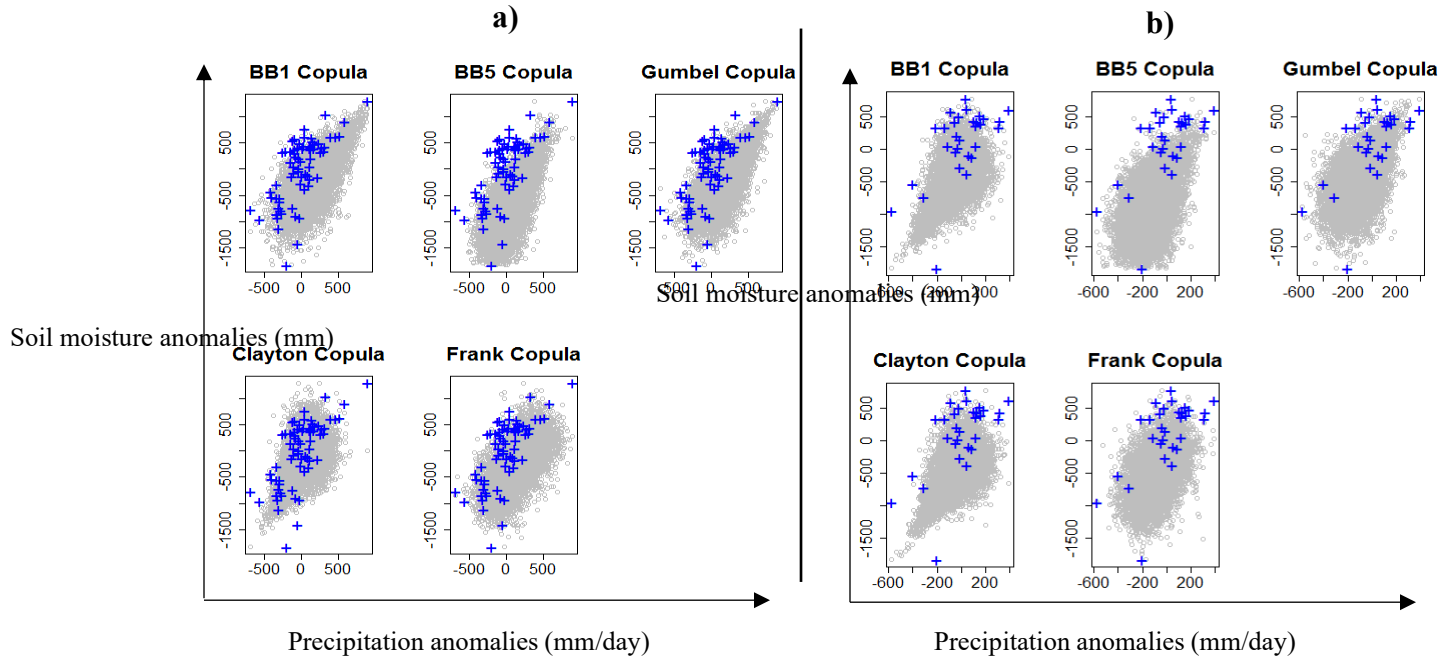


Figure 4.5. Scatter plot of observed data and 10 000 generated pairs from the fitted copula, for modeling dependence structure between precipitation and soil moisture anomalies in northern California, USA: a) from 1948 to 2017 and b) from 1988 to 2017

It is worth noting that the two data sets have relatively small dependence strength. The Kendall's τ is 0.48 and 0.35, for data set of length 70 and 30 years, respectively. Copulas with dependence strength not matching with those of the series are discarded. Five copulas are retained for subsequent analysis: M-copulas (*BB1* and *BB5*) and one-parameter copulas (Gumbel, Frank and Clayton).

From Figure 4.5a, in the case $n = 70$, a visual inspection of different copulas fitted to the data shows that only the Clayton and Frank copulas are able to characterize the dependence structure of this data set. However, from Figure 4.6a, theoretical Kendall's function of the *BB5* and Clayton copulas are closer to the empirical Kendall's function than other copulas. Then, *BB5* and Clayton copulas could better model the dependence between precipitation and soil moisture anomalies, when the original data record is considered ($n=70$).

Results regarding the application of the proposed *GOF* tests as well as classical tests are summarized in Table 4.10. For the 70-years data set, neither Clayton, Frank nor Gumbel copulas could be rejected, considering the *p-values*, by Z_n^{LM} and S_n^C .

Table 4.10. *P-values* for the *GOF* tests for modeling dependence structure of precipitation and soil moisture anomalies in northern California, USA

Sample size	Copula family	Copula parameters		Proposed <i>GOF</i> tests (<i>p-values</i>)		Classical tests (<i>p-values</i>)		<i>AIC</i>
		θ_1	θ_2	Z_n^X	Z_n^{LM}	S_n^C	S_n^K	
<i>n</i> =30	<i>BB1</i>	1.21	1.01	0.04	0.01	0.05	0.03	8.53
	<i>BB5</i>	0.89	0.94	0.48	0.13	0.04	0.47	3.84
	Gumbel	1.52	-	0.25	0.32	0.01	0.47	5.48
	Clayton	1.23	-	0.04	0.03	0.04	0.46	10.5
	Frank	3.61	-	0.13	0.07	0.01	0.15	6.75
<i>n</i> =70	<i>BB1</i>	0.001	2.02	0.04	0.07	0.06	0.05	47.4
	<i>BB5</i>	0.88	1.27	0.16	0.10	0.05	0.02	28.1
	Gumbel	2.02	-	0.04	0.17	0.87	0.02	48.4
	Clayton	1.16	-	0.46	0.89	0.25	0.23	27.4
	Frank	5.54	-	0.03	0.48	0.85	0.06	46.06

This table presents the estimated *p-values* at the significance level $\alpha = 5\%$ for each test. Values in bold indicate the copulas that pass the *GOF* tests. The grey color indicates the best fitting copula regarding the model selection criterion (*AIC*).

It is also notable, from Table 4.10, that the proposed statistic Z_n^X , accept only the *BB5* and Clayton copulas as good candidates to fit data. It is also important to highlight that the *BB5* copula is rejected by the classical statistics S_n^K and S_n^C . Recall that *BB5* and Clayton copulas provide a very good fit to the data regarding the Kendall plot (Figure 4.6a).

Since *p-values* cannot be used to rank copulas, but only to accept or reject them (e.g. Salvadori and Michele, 2011), we used the criterion *AIC* for ranking copula that passed *GOF* tests. Consequently, both Clayton and *BB5* copulas could be chosen as the best fitting copula, since they have very closer *AIC*.

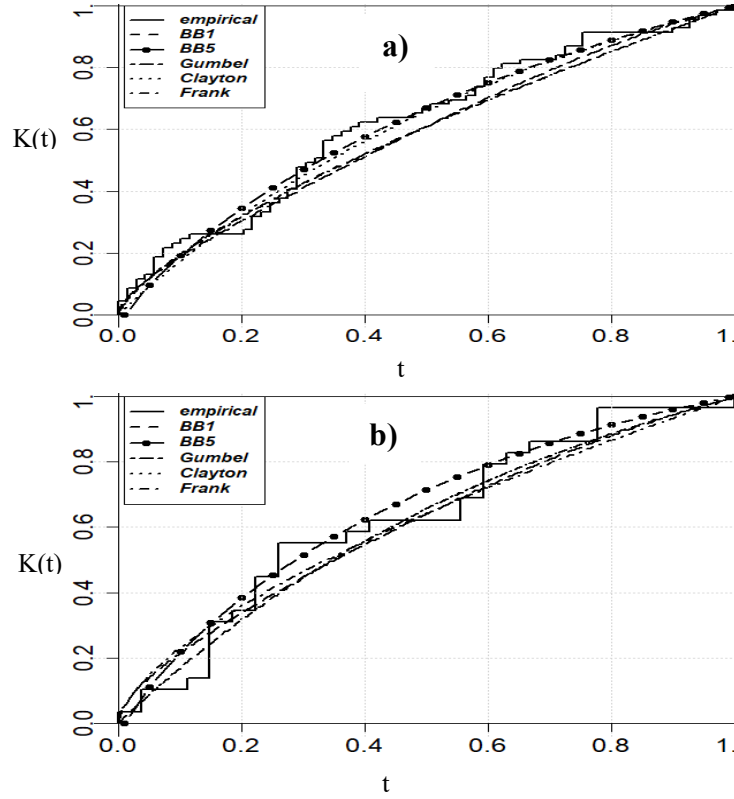


Figure 4.6. Comparison between the empirical estimate and theoretical Kendall's function for each fitted copula for modeling the dependence structure between precipitation and soil moisture anomalies in northern California, USA: a) from 1948 to 2017 and b) from 1988 to 2017

In the case of small sample size, the graphical diagnostics of the dependence structure (Figures 4.5b and 4.6b), in order to select the appropriate copulas, become difficult. From Figure 4.5b, we can see that all fitted copula could be a good candidate. This behavior is replicated in Figure 4.6b where different theoretical and empirical Kendall's functions are close to each other. This highlights the importance of using formal *GOF* tests. When dealing with small sample size ($n=30$), based on Z_n^X and Z_n^{LM} statistics, *BB5*, Gumbel and Frank copulas seem to fit properly the data (Table 4.10). Based on S_n^K , all considered copulas are able to represent the dependence structure, except the *BB1* copula, whereas none of these copulas is accepted by S_n^C . This might be explained by the inconsistency of these two tests for small sized sample (Genest *et al.*, 2006). According to *AIC* criterion, *BB5* copula is selected as the best copula followed by Gumbel copula.

6. Conclusion

The proposed *GOF* tests are the first ones designed specifically to deal with *M*-copulas. A novelty of this paper lies in the introduction of multivariate L-moments in *GOF* copula testing. Simulations results show that the proposed tests are good candidates to control first type errors and overall, they have the highest performance in terms of power (even for small samples and weak dependence). Furthermore, the obtained results from two different case studies confirm the advantages of the proposed tests especially when dealing with small sample sizes as well as when considering several Extreme-Value copulas.

Troisième Partie

Chapitre 5. Conclusions et perspectives

Ce chapitre conclut la thèse en résumant les principaux résultats obtenus au cours de ces travaux. Nous discutons également des perspectives ouvertes qu'offrent ces résultats.

1. Conclusions générales

Les travaux présentés dans cette thèse s'inscrivent dans le cadre général de l'*AFM* des variables hydrologiques. En effet, l'analyse et la modélisation adéquate des variables hydrologiques sont cruciales pour l'étude des événements extrêmes, telle que la crue. Dans ce cadre, l'*AF* est un des outils privilégiés pour l'estimation des quantiles associés aux extrêmes hydrologiques. Par ailleurs, généralement, ces événements extrêmes sont caractérisés par plusieurs variables dépendantes. Par exemple, une crue peut se définir en fonction de sa pointe, sa durée et son volume. Par conséquent, les deux dernières décennies ont vu un grand développement de la modélisation de la dépendance de ces variables, en se basant principalement sur les copules. Cependant, les observations doivent vérifier les hypothèses d'homogénéité, stationnarité et indépendance. Or, ces hypothèses ne sont pas souvent respectées, notamment en ce qui concerne l'homogénéité des observations. De plus, la majorité des études existantes se concentre sur l'étape de modélisation en omettant les autres étapes de l'*AFM*. D'où la nécessité d'étudier l'*AFM* et ces différentes étapes dans un cadre d'absence de l'hypothèse d'homogénéité. Le principal objectif des méthodologies proposées dans cette thèse étant de remédier aux inconvénients des méthodes classiques d'*AFM* pour améliorer la qualité de l'estimation du risque associé aux événements extrêmes. Particulièrement, les nouvelles approches proposées visent à traiter l'homogénéité des séries hydrologiques qui touchent de près à la modélisation des crues en *AFM*. D'abord, nous avons proposé un test de détection de l'hétérogénéité des séries hydrologiques multivariées. Ensuite, nous avons introduit un modèle de

mélange de copules pour la modélisation de l'hétérogénéité. Finalement, nous avons proposé des tests d'adéquation pour les copules multiparamètres.

Dans un premier temps, l'accent a été mis sur la vérification de la validité de l'hypothèse de l'homogénéité multivariée. Ceci a fait l'objet de l'article « *Homogeneity testing of multivariate hydrological records, using multivariate copula L-moments* ». Motivés par l'insuffisance des méthodes employées actuellement en hydrologie, nous nous sommes plus particulièrement concentrés sur les aspects méthodologiques du traitement statistique de l'hétérogénéité. Dans le cadre de cette étape, nous avons proposé un nouveau test capable de détecter la rupture dans la structure de dépendance décrite par la copule. La statistique du test proposé est basée sur les L-moments multivariés. Ce nouveau test a permis de relaxer les hypothèses trop rigides imposées par les tests existants. Particulièrement, le nouveau test permet aux marges d'être hétérogènes. Ainsi, il se présente comme un moyen efficace et robuste pour déceler une rupture dans la dépendance, c'est-à-dire de tester l'existence d'un changement à une position donnée aussi bien dans la mesure de dépendance que dans le type de la copule. La statistique de test ainsi définie est non paramétrique et valable pour tester les séries de petite taille. Par ailleurs, elle ne nécessite de faire aucune hypothèse sur la structure de dépendance des données, ce qui la rend plus flexible et adaptée aux séries hydrologiques et à l'AF.

La performance du test proposé est évaluée à travers une étude de simulations exhaustive. Ainsi, dans le but d'étudier les propriétés du test (erreur de première espèce et puissance), nous avons utilisé diverses combinaisons de types de rupture, d'emplacements de rupture et de tailles d'échantillon. La technique du bootstrap a été utilisée pour obtenir les *p-values*. Les simulations ont montré clairement que l'utilisation de statistiques sommaires, comme les L-moments, est une alternative avantageuse aux tests existants pour détecter la rupture dans la force et/ou le type de la

dépendance. En effet, le test proposé contrôle bien l'erreur de type I et il a généralement une bonne puissance. De plus, les puissances tendent à augmenter en fonction de la taille d'échantillon. Tel qu'attendu, plus l'écart dans la force de dépendance (tel que décrite par le tau de Kendall) est élevé, meilleures sont les puissances, peu importe le lieu de la rupture. La bonne qualité des résultats obtenus dans la présente étude s'explique par le fait que les L-moments détectent les caractéristiques qui confèrent aux copules leurs signatures pour décrire la structure de dépendance.

Dans le cadre d'une rupture dans la force de dépendance seulement, nous avons fait une comparaison, sur la base de simulations, entre le test proposé et ceux qui sont couramment utilisés, notamment le test basé sur le processus de Kendall et celui basé sur le rapport de vraisemblance. Ce dernier a une faible puissance de détection des ruptures. Toutefois, le test de Kendall s'est montré très prometteur pour détecter les hétérogénéités. Ce test offre néanmoins les meilleures puissances de détection lorsque la variation du tau de Kendall est importante. Pour de faibles amplitudes de changement, les deux tests classiques ne détectent pas le point de rupture, peu importe la taille de l'échantillon. Par ailleurs, cette comparaison montre que le test proposé peut être bénéfique dans certaines situations. En effet, la nouvelle statistique est plus flexible que la plupart des techniques existantes dans le cas où il n'y a pas de variation importante dans le tau de Kendall ou aussi pour les échantillons de petite taille. De plus, le test proposé a une meilleure performance dans la localisation des points de changements.

La validation du test proposé est illustrée plus spécifiquement par une application sur deux cas d'études. Ces applications permettent de montrer comment interpréter les résultats du nouveau test et quels en sont les avantages. La première illustration correspond à une série hydrologique de taille moyenne et pour laquelle seulement le type de la dépendance change. La deuxième application traite une série de grande taille et caractérisée par un changement à la fois dans le type et la force

de la dépendance. L'interprétation des résultats permet de mettre en exergue les avantages du test proposé dans la détection des différentes ruptures dans la structure de dépendance. En effet, dans les deux cas, le test proposé a permis de détecter le changement et de bien localiser le moment de la rupture dans la dépendance.

Dans un deuxième temps, nous nous sommes penchés sur l'étape de modélisation en l'absence de l'hypothèse d'homogénéité. Cette partie a fait l'objet du deuxième article « *Meta-heuristic estimation method for mixture copula models* ». Cette étude est motivée par le constat que ne pas considérer la possibilité qu'une rupture soit survenue dans l'historique des débits peut être lourd de conséquences. En effet, si les observations ont un changement dans leur structure de dépendance, les estimations des quantiles reposant sur l'hypothèse d'homogénéité engendreront des conclusions erronées sur l'évènement extrême étudié. Elles peuvent, par exemple, conduire au sous-dimensionnement ou surdimensionnement des nouveaux ouvrages hydrauliques et entraîner ainsi des pertes financières importantes et/ou des pertes de vies humaines. Toutefois, les approches classiques de modélisation faisant appel à l'utilisation des copules ignorent l'hétérogénéité des données. Nous avons donc proposé l'utilisation d'un modèle de mélange de copules (copules mixtes) afin de tenir compte de l'hétérogénéité de la structure de dépendance pour la modélisation des séries hydrologiques multivariées. Deux types de modèles sont définis: homogène et hétérogène. Le modèle de mélange est décrit comme homogène si les deux composantes appartiennent à la même famille. Dans le cas contraire, le modèle mixte est considéré comme hétérogène.

Dans le cadre de cette thèse, les paramètres du modèle de copules mixtes ainsi proposé sont estimés par la méthode du maximum de pseudo-vraisemblance. Or, la maximisation de cette fonction n'est pas directe comme dans le cas du modèle simple. En effet, le modèle mixte compte de nombreux

paramètres à estimer, et les méthodes classiques d'optimisation sont inefficaces dans ce contexte. Par conséquent, nous avons proposé une approche de maximisation basée sur les algorithmes génétiques. Ces algorithmes permettent la résolution de problèmes d'optimisation complexes tout en respectant les caractéristiques des paramètres. La méthode proposée est dénotée *MH-MPL* (Meta-Heuristic Maximum Pseudo-Likelihood).

Nous avons dans un second temps utilisé l'approche Expectation-Maximisation (*EM*) comme outil d'optimisation de la fonction de pseudo-vraisemblance. Les performances de deux approches d'estimation, dans un contexte hydrologique, sont évaluées par une étude de simulations. Par ailleurs, nous avons effectué une étude de sensibilité aux différents facteurs susceptibles d'influencer une méthode d'estimation, particulièrement la taille d'échantillons, le ratio de mélange et le type ainsi que la force de la dépendance. Les résultats de simulations montrent que, globalement, la performance des deux approches s'améliore quand la taille d'échantillon augmente et aussi lorsque la force de dépendance est plus forte. En effet, les critères de performance *RRMSE* et *RBIAS* sont plus faibles pour les tailles d'échantillons et mesures de dépendance les plus grandes. De plus, nous avons remarqué que les deux estimateurs sont plus robustes quand le ratio de mélange est égal à 0.5. Toutefois, bien que la performance des deux méthodes évolue de la même façon sous différents scénarios, l'étude comparative montre que l'utilisation des algorithmes génétiques est plus avantageuse pour la maximisation de la pseudo-vraisemblance. En fait, la méthode *MH-MPL* permet d'obtenir de meilleurs résultats au niveau du *RRMSE* et du *RBIAS*. De plus, la méthode *MH-MPL* s'est montrée supérieure en améliorant la précision d'estimation des paramètres du modèle mixte. En particulier, il a été montré que la variance d'estimation, associée à la méthode *MH-MPL*, a été considérablement réduite comparée à la méthode *EM*. L'amplitude de cette

réduction est plus importante principalement dans le cas de données faiblement dépendantes ou aussi de faible taille, ce qui est intéressant pour les applications hydrologiques.

Une fois le modèle fréquentiel choisi et ses paramètres estimés, il est primordial de vérifier son adéquation à la série étudiée. Cette partie a fait l'objet du troisième article « *Multivariate L-moment based tests for copula selection, with hydrometeorological applications* ». Malgré l'expansion de l'application des copules en *AFM*, il subsiste encore des déficiences plus ou moins sérieuses telles que, entre autres, l'absence des procédures d'adéquation spécifiques pour les copules multiparamètres. Dans le cadre de la présente thèse, nous avons proposé deux tests d'adéquation pour les copules multiparamètres. Les nouveaux tests sont basés sur les L-moments multivariés. Ces tests permettent de mieux vérifier l'adéquation des familles de copules qu'avec les tests classiques. La technique du bootstrap a été employée pour calculer les *p-values* des tests proposés. Nous avons mené des simulations pour évaluer la performance des tests développés et les comparer avec ceux existants. Les résultats ont permis de constater que les nouveaux tests sont généralement efficaces sous divers scénarios de dépendance. Aussi, la puissance des tests augmente avec la taille de l'échantillon et la force de la dépendance. Les tests ont néanmoins les meilleures puissances lorsque le tau de Kendall est plus élevé. Par ailleurs, il a été également constaté que les nouveaux tests se comportent mieux que ceux existants. En particulier, les nouvelles statistiques sont plus adaptées pour vérifier l'adéquation des copules de type valeurs extrêmes. De plus, ces tests conservent leur seuil nominal même pour les échantillons de petite taille. Ainsi, ces derniers s'adaptent mieux aux spécificités des séries hydrologiques. Enfin, les deux tests ont été mis en œuvre sur des données réelles. D'abord, deux séries de données correspondant aux variables décrivant la crue sont utilisées. La première série est caractérisée par un faible tau de Kendall alors que la deuxième est caractérisée par un tau de Kendall élevé. Ceci nous a permis de vérifier l'effet

de tau de Kendall sur la performance des tests proposés. Ensuite, nous avons appliqué les tests sur des séries de données correspondant aux variables caractérisant le phénomène de la sécheresse. Dans ce cas, nous nous sommes intéressés à deux séries de différentes tailles. Ces applications ont montré, d'une part, l'importance de la vérification de l'adéquation et, d'autre part, l'avantage des nouveaux tests dans la validation du modèle ajusté.

2. Perspectives

Certaines interrogations soulevées tout au long du développement et de l'application des méthodes présentées dans le cadre de la présente thèse ouvrent des perspectives intéressantes pour des travaux futurs. Les principales perspectives de recherche qui peuvent être envisagées à l'issue de ces études sont :

1. Détection des ruptures multiples et graduelles: Dans la même direction que le test de détection de rupture développé dans le cadre de cette thèse (Chapitre 2), un test pour la détection simultanée des ruptures multiples mérite une investigation. Également, il serait sans doute intéressant de fonder, en utilisant les L-moments, un test de rupture graduelle dans la structure de dépendance. Celui-ci généraliserait le test d'un changement abrupt de dépendance, qui est développé dans cette thèse. Dans le cadre de ce travail, seul le cas bivarié a été étudié par simulations. Il serait également intéressant d'explorer davantage le cas de plusieurs variables.

2. Estimation des paramètres du modèle de copules mixtes : L'approche proposée dans le cadre de cette thèse est basée sur les algorithmes génétiques (*AG*). Une direction future consiste à utiliser et comparer d'autres types d'algorithmes méta-heuristiques dont font partie les *AG*. De plus, il est envisageable d'adapter la méthode *EM* afin d'améliorer la qualité des estimations. Dans ce contexte, il est possible d'utiliser d'autres variantes de l'Algorithme *EM* telles que l'algorithme du

gradient *EM*, l'algorithme *Stochastic EM* ou encore l'algorithme *Monte-Carlo EM*. De plus, il serait intéressant d'évaluer les incertitudes associées aux différents estimateurs.

3. Test d'adéquation pour les copules : La performance des approches proposées a été évaluée uniquement dans un contexte hydrologique. Il serait intéressant d'étudier cette performance dans d'autres contextes tels que, par exemple, la finance ou l'actuariat. Il serait également important de développer un test d'adéquation qui prend en considération la non-stationnarité des séries hydrologiques multivariées. Le cas échéant, ce test permettrait la sélection du meilleur modèle en présence de plusieurs covariables afin de mieux estimer le risque des extrêmes hydrologiques tout en tenant compte des aléas climatiques.

4. Modèle de copules mixtes non stationnaires : Le modèle mixte proposé dans le cadre de cette thèse est un modèle stationnaire dans le sens où chaque composante est stationnaire. Une autre avenue envisageable concerne le développement d'un modèle à rupture de dépendance non stationnaire. Ce type de modèles permettrait aux composantes de varier en fonction du temps ou de covariables climatiques.

5. Estimation des quantiles en absence d'homogénéité : Pour un risque p , le quantile bivarié estimé par l'*AFM* est représenté par une courbe avec un nombre infini des couples (de pointes Q et de volumes V dans le cadre de la crue). En absence d'homogénéité, le comportement et la forme de ces courbes n'ont pas été étudiés dans le cadre de la présente thèse. Il serait intéressant d'examiner l'impact de l'hétérogénéité sur ces courbes dans les travaux futurs. De plus, le choix du couple optimal qui peut être considéré dans un cas pratique n'a pas été l'objet de travaux approfondis. Une étude approfondie qui permettrait de déterminer le ou les couples optimaux à considérer lors d'une étude pratique est envisageable.

Quatrième partie

Chapitre 6. Références bibliographiques

- Adamowski K (1985) Nonparametric kernel estimation of flood frequencies. *Water Resources Research* 21(11):1585-1590.
- AghaKouchak A (2014) Entropy-copula in hydrology and climatology. *Journal of Hydrometeorology* 15(6):2176-2189.
- Ailliot P, Thompson C and Thomson P (2009) Space-time modelling of precipitation by using a hidden Markov model and censored Gaussian distributions. *Journal of the Royal Statistical Society: Series C (Applied Statistics)* 58(3):405-426.
- Aissia MAB, Chebana F, Ouarda TB, Roy L, Desrochers G, Chartier I and Robichaud É (2012) Multivariate analysis of flood characteristics in a climate change context of the watershed of the Baskatong reservoir, Province of Québec, Canada. *Hydrological Processes* 26(1):130-142.
- Akaike H (1974) A new look at the statistical model identification. *Selected Papers of Hirotugu Akaike*, Springer. p 215-222.
- Alila Y and Mtiraoui A (2002) Implications of heterogeneous flood-frequency distributions on traditional stream-discharge prediction techniques. *Hydrological Processes* 16(5):1065-1084.
- Arakelian V and Karlis D (2014) Clustering dependencies via mixtures of copulas. *Communications in Statistics-Simulation and Computation* 43(7):1644-1661.
- Arcidiacono P and Jones JB (2003) Finite mixture distributions, sequential likelihood and the EM algorithm. *Econometrica* 71(3):933-946.
- Ashkar F and Aucoin F (2011) A broader look at bivariate distributions applicable in hydrology. *Journal of Hydrology* 405(3-4):451-461.
- Barbe P, Genest C, Ghouli K and Rémillard B (1996) On Kendall's process. *Journal of Multivariate Analysis* 58(2):197-229.
- Barber C, Lamontagne JR and Vogel RM (2020) Improved estimators of correlation and R2 for skewed hydrologic data. *Hydrological Sciences Journal* 65(1):87-101.
- Barth NA, Villarini G, Nayak MA and White K (2017) Mixed populations and annual flood frequency estimates in the western United States: The role of atmospheric rivers. *Water Resources Research* 53(1):257-269.
- Beaulieu C, Seidou O, Ouarda TB and Zhang X (2009) Intercomparison of homogenization techniques for precipitation data continued: Comparison of two recent Bayesian change point models. *Water Resources Research* 45(8).
- Ben Nasr I and Chebana F (2019a) Homogeneity testing of multivariate hydrological records, using multivariate copula L-moments. *Advances in Water Resources* 134:103449.
- Ben Nasr I and Chebana F (2019b) Multivariate L-moment based tests for copula selection, with hydrometeorological applications. *Journal of Hydrology* 579:124151.
- Benamer S, Benkhaled A, Meraghni D, Chebana F and Necir A (2017) Complete flood frequency analysis in Abiod watershed, Biskra (Algeria). *Natural Hazards* 86(2):519-534.
- Bender J, Wahl T and Jensen J (2014) Multivariate design in the presence of non-stationarity. *Journal of Hydrology* 514:123-130.
- Berg D (2009) Copula goodness-of-fit testing: an overview and power comparison. *The European Journal of Finance* 15(7-8):675-701.
- Berg D and Quessy JF (2009) Local Power Analyses of Goodness-of-fit Tests for Copulas. *Scandinavian Journal of Statistics* 36(3):389-412.
- Bilgrau AE, Eriksen PS, Rasmussen JG, Johnsen HE, Dybkær K and Bøgsted M (2016) GMCM: Unsupervised clustering and meta-analysis using gaussian mixture copula models. *Journal of Statistical Software* 70(2):1-23.

- Bonanomi A, Nai Ruscone M and Osmetti SA (2019) Dissimilarity measure for ranking data via mixture of copulae. *Statistical Analysis and Data Mining: The ASA Data Science Journal*.
- Bouzebda S and Keziou A (2013) A semiparametric maximum likelihood ratio test for the change point in copula models. *Statistical Methodology* 14:39-61.
- Brahimi B, Chebana F and Necir A (2015) Copula representation of bivariate L-moments: a new estimation method for multiparameter two-dimensional copula models. *Statistics* 49(3):497-521.
- Brahimi B and Necir A (2012) A semiparametric estimation of copula models based on the method of moments. *Statistical Methodology* 9(4):467-477.
- Bryden E, Carlson JB and Craig B (1995) *Some Monte Carlo results on nonparametric changepoint tests*. Citeseer,
- Bücher A and Volgushev S (2013) Empirical and sequential empirical copula processes under serial dependence. *Journal of Multivariate Analysis* 119:61-70.
- Buishand TA, De Martino G, Spreeuw J and Brandsma T (2013) Homogeneity of precipitation series in the Netherlands and their trends in the past century. *International journal of climatology* 33(4):815-833.
- Burn DH and Elnur MAH (2002) Detection of hydrologic trends and variability. *Journal of hydrology* 255(1):107-122.
- Calenda G, Mancini CP and Volpi E (2009) Selection of the probabilistic model of extreme floods: The case of the River Tiber in Rome. *Journal of Hydrology* 371(1-4):1-11.
- Capérea P, Fougères A-L and Genest C (1997) A nonparametric estimation procedure for bivariate extreme value copulas. *Biometrika* 84(3):567-577.
- Capérea P and Genest C (1993) Spearman's ρ is larger than Kendall's τ for positively dependent random variables. *Journal of Nonparametric Statistics* 2(2):183-194.
- Caudill SB and Acharya RN (1998) Maximum likelihood estimation of a mixture of normal regressions: starting values and singularities. *Communications in Statistics-Simulation and Computation* 27(3):667-674.
- Celeux G, Chrétien S, Forbes F and Mkhadri A (2001) A component-wise EM algorithm for mixtures. *Journal of Computational and Graphical Statistics* 10(4):697-712.
- Chang R and Wang M (1983) Shifted Legendre direct method for variational problems. *Journal of Optimization Theory and Applications* 39(2):299-307.
- Chebana F (2013) Multivariate analysis of hydrological variables. *Encyclopedia of Environmetrics*.
- Chebana F, Aissia M-AB and Ouarda TB (2017) Multivariate shift testing for hydrological variables, review, comparison and application. *Journal of Hydrology* 548:88-103.
- Chebana F, Ben Aissia M-A and Ouarda BMJT (2016) Multivariate shift testing for hydrological variables, review, comparison and application. *journal of hydrology*.
- Chebana F and Ouarda TB (2007) Multivariate L-moment homogeneity test. *Water resources research* 43(8).
- Chebana F and Ouarda TB (2009) Index flood-based multivariate regional frequency analysis. *Water Resources Research* 45(10).
- Chebana F and Ouarda TB (2011) Multivariate quantiles in hydrological frequency analysis. *Environmetrics* 22(1):63-78.
- Chebana F, Ouarda TB and Duong TC (2013) Testing for multivariate trends in hydrologic frequency analysis. *Journal of hydrology* 486:519-530.
- Chebana F, Ouarda TBMJ and Duong TC (2010) Testing for stationarity, homogeneity and trend in multivariate hydrologic time series: a review. (INRS-ETE, Canada Research Chair on the Estimation of Hydrometeorological Variables), p 71.

- Chen K-H and Khashanah K (2014) Measuring systemic risk: copula CoVaR. *Available at SSRN* 2473648.
- Choulakian V and Stephens MA (2001) Goodness-of-fit tests for the generalized Pareto distribution. *Technometrics* 43(4):478-484.
- Chowdhury JU, Stedinger JR and Lu LH (1991) Goodness-of-fit tests for regional generalized extreme value flood distributions. *Water Resources Research* 27(7):1765-1776.
- Christensen TS, Pircalabu A and Høg E (2019) A seasonal copula mixture for hedging the clean spark spread with wind power futures. *Energy Economics* 78:64-80.
- Das J and Umamahesh N (2017) Uncertainty and Nonstationarity in Streamflow Extremes under Climate Change Scenarios over a River Basin. *Journal of Hydrologic Engineering* 22(10):04017042.
- De Michele C, Salvadori G, Canossi M, Petaccia A and Rosso R (2005) Bivariate statistical approach to check adequacy of dam spillway. *Journal of Hydrologic Engineering* 10(1):50-57.
- De Michele C, Salvadori G, Vezzoli R and Pecora S (2013) Multivariate assessment of droughts: Frequency analysis and dynamic return period. *Water Resources Research* 49(10):6985-6994.
- Dehling H, Vogel D, Wendler M and Wied D (2017) Testing for changes in Kendall's tau. *Econometric Theory* 33(6):1352-1386.
- Dias A and Embrechts P (2004) Change-point analysis for dependence structures in finance and insurance. *Risk Measures of the 21st Century* :321335.
- Ding W and Song PX-K (2016) EM algorithm in Gaussian copula with missing data. *Computational Statistics & Data Analysis* 101:1-11.
- Dou X, Kuriki S, Lin GD and Richards D (2016) EM algorithms for estimating the Bernstein copula. *Computational Statistics & Data Analysis* 93:228-245.
- Durocher M, Chebana F and Ouarda TB (2015) A Nonlinear Approach to Regional Flood Frequency Analysis Using Projection Pursuit Regression. *Journal of Hydrometeorology* 16(4):1561-1574.
- Durocher M, Chebana F and Ouarda TB (2016) On the prediction of extreme flood quantiles at ungauged locations with spatial copula. *Journal of Hydrology* 533:523-532.
- Durocher M and Quessy JF (2017) Goodness-of-fit tests for copula-based spatial models. *Environmetrics* 28(5):e2445.
- Ehsanzadeh E, Ouarda TB and Saley HM (2011) A simultaneous analysis of gradual and abrupt changes in Canadian low streamflows. *Hydrological Processes* 25(5):727-739.
- El Adlouni S, Ouarda T, Zhang X, Roy R and Bobée B (2007) Generalized maximum likelihood estimators for the nonstationary generalized extreme value model. *Water Resources Research* 43(3).
- Everitt BS (2014) Finite mixture distributions. *Wiley StatsRef: Statistics Reference Online*.
- Evin G, Merleau J and Perreault L (2011) Two-component mixtures of normal, gamma, and Gumbel distributions for hydrological applications. *Water Resources Research* 47(8).
- Fan Y, de Micheaux PL, Penev S and Salopek D (2017) Multivariate nonparametric test of independence. *Journal of Multivariate Analysis* 153:189-210.
- Fan Y, Huang W, Huang G, Huang K, Li Y and Kong X (2015) Bivariate hydrologic risk analysis based on a coupled entropy-copula method for the Xiangxi River in the Three Gorges Reservoir area, China. *Theoretical and Applied Climatology* :1-17.

- Fan Y, Huang W, Huang G, Li Y, Huang K and Li Z (2016) Hydrologic risk analysis in the Yangtze River basin through coupling Gaussian mixtures into copulas. *Advances in water resources* 88:170-185.
- Fermanian J-D (2005) Goodness-of-fit tests for copulas. *Journal of multivariate analysis* 95(1):119-152.
- Fermanian J-D (2013) An overview of the goodness-of-fit test problem for copulas. *Copulae in Mathematical and Quantitative Finance*, Springer. p 61-89.
- Franks SW and Kuczera G (2002) Flood frequency analysis: Evidence and implications of secular climate variability, New South Wales. *Water resources research* 38(5).
- Fu G and Butler D (2014) Copula-based frequency analysis of overflow and flooding in urban drainage systems. *Journal of Hydrology* 510:49-58.
- Fu Y, Liu Y, Wang H-H and Wang X (2019) Empirical likelihood estimation in multivariate mixture models with repeated measurements. *Statistical Theory and Related Fields* :1-9.
- Genest C and Boies J-C (2003) Detecting dependence with Kendall plots. *The American statistician* 57(4):275-284.
- Genest C and Chebana F (2016) Copula modeling in hydrologic frequency analysis. *Handbook of Applied Hydrology* V.P. Singh (Édit.)McGraw-Hill, New York, In press Ed.
- Genest C and Chebana F (2017) Copula modeling in hydrological frequency analysis. *Handbook of Applied Hydrology, Second Edition (Chapter 30)*, Singh VP (Édit.)McGraw-Hill, New York. p 301-309.
- Genest C and Favre A-C (2007) Everything you always wanted to know about copula modeling but were afraid to ask. *Journal of hydrologic engineering* 12(4):347-368.
- Genest C, Ghouli K and Rivest L-P (1995) A semiparametric estimation procedure of dependence parameters in multivariate families of distributions. *Biometrika* 82(3):543-552.
- Genest C, Kojadinovic I, Nešlehová J and Yan J (2011a) A goodness-of-fit test for bivariate extreme-value copulas. *Bernoulli* 17(1):253-275.
- Genest C and Nešlehová J (2014) Copulas and copula models. *Wiley StatsRef: Statistics Reference Online*.
- Genest C, Nešlehová J and Ben Ghorbal N (2011b) Estimators based on Kendall's tau in multivariate copula models. *Australian & New Zealand Journal of Statistics* 53(2):157-177.
- Genest C, Quessy J-F and Remillard B (2006) Goodness-of-fit procedures for copula models based on the probability integral transformation. *Scandinavian Journal of Statistics* :337-366.
- Genest C and Remillard B (2008) Validity of the parametric bootstrap for goodness-of-fit testing in semiparametric models. *Annales de l'IHP Probabilités et statistiques*. p 1096-1127.
- Genest C, Remillard B and Beaudoin D (2009) Goodness-of-fit tests for copulas: A review and a power study. *Insurance: Mathematics and economics* 44(2):199-213.
- Genest C and Rivest L-P (1993) Statistical inference procedures for bivariate Archimedean copulas. *Journal of the American statistical Association* 88(423):1034-1043.
- Gilroy KL and McCuen RH (2012) A nonstationary flood frequency analysis method to adjust for future climate change and urbanization. *Journal of Hydrology* 414:40-48.
- Gombay E and Horvath L (1999) Change-points and bootstrap. *Environmetrics* 10(6):725-736.
- Gómez M, Ausín MC and Domínguez MC (2017) Seasonal copula models for the analysis of glacier discharge at King George Island, Antarctica. *Stochastic Environmental Research and Risk Assessment* 31(5):1107-1121.
- Good PI (2004) *Permutation, Parametric, and Bootstrap Tests of Hypotheses (Springer Series in Statistics)*. Springer-Verlag New York, Inc.,

- Grego JM and Yates PA (2010) Point and standard error estimation for quantiles of mixed flood distributions. *Journal of hydrology* 391(3-4):289-301.
- Grimaldi S and Serinaldi F (2006) Asymmetric copula in multivariate flood frequency analysis. *Advances in Water Resources* 29(8):1155-1167.
- Gudendorf G and Segers J (2010) Extreme-value copulas. *Copula theory and its applications*, Springer. p 127-145.
- Guegan D and Zhang J (2010) Change analysis of a dynamic copula for measuring dependence in multivariate financial data. *Quantitative Finance* 10(4):421-430.
- Guerfi N, Assani AA, Mesfioui M and Kinnard C (2015) Comparison of the temporal variability of winter daily extreme temperatures and precipitations in southern Quebec (Canada) using the Lombard and copula methods. *International Journal of Climatology* 35(14):4237-4246.
- Hamed K and Rao AR (1999) *Flood frequency analysis*. CRC press,
- Han J-C, Huang G-H, Zhang H, Li Z and Li Y-P (2014) Bayesian uncertainty analysis in hydrological modeling associated with watershed subdivision level: a case study of SLURP model applied to the Xiangxi River watershed, China. *Stochastic environmental research and risk assessment* 28(4):973-989.
- Hao Z and AghaKouchak A (2013) Multivariate standardized drought index: a parametric multi-index model. *Advances in Water Resources* 57:12-18.
- Hao Z and Singh VP (2016) Review of dependence modeling in hydrology and water resources. *Progress in Physical Geography* 40(4):549-578.
- Hassanzadeh Y, Abdi A, Talatahari S and Singh VP (2011) Meta-heuristic algorithms for hydrologic frequency analysis. *Water Resources Management* 25(7):1855-1879.
- He W-P, Liu Q-Q, Gu B and Zhao S-S (2016) A novel method for detecting abrupt dynamic change based on the changing Hurst exponent of spatial images. *Climate dynamics* 47(7-8):2561-2571.
- Hoff PD (2007) Extending the rank likelihood for semiparametric copula estimation. *The Annals of Applied Statistics* 1(1):265-283.
- Holmes M, Kojadinovic I and Quessy J-F (2013) Nonparametric tests for change-point detection à la Gombay and Horváth. *Journal of Multivariate Analysis* 115:16-32.
- Hosking J and Wallis J (1993) Some statistics useful in regional frequency analysis. *Water Resources Research* 29(2):271-281.
- Hosking JRM (1990) L-Moments: Analysis and Estimation of Distributions Using Linear Combinations of Order Statistics. *Journal of the Royal Statistical Society. Series B (Methodological)* 52(1):105-124.
- Hu L (2006) Dependence patterns across financial markets: a mixed copula approach. *Applied financial economics* 16(10):717-729.
- Hundecha Y, Pahlow M and Schumann A (2009) Modeling of daily precipitation at multiple locations using a mixture of distributions to characterize the extremes. *Water resources research* 45(12).
- Jiang C, Xiong L, Xu CY and Guo S (2015) Bivariate frequency analysis of nonstationary low-flow series based on the time-varying copula. *Hydrological Processes* 29(6):1521-1534.
- Joe H (2014) *Dependence modeling with copulas*. CRC Press,
- Joe H and Xu JJ (1996) The estimation method of inference functions for margins for multivariate models.
- John G (2014) Encyclopedia of Environmetrics (2nd edition). *Reference Reviews* 28(1):27-28.
- Joshi D, St-Hilaire A, Ouarda T and Daigle A (2015) Statistical downscaling of precipitation and temperature using sparse Bayesian learning, multiple linear regression and genetic

- programming frameworks. *Canadian Water Resources Journal/Revue Canadienne Des Ressources Hydriques* 40(4):392-408.
- Kao S-C and Govindaraju RS (2010) A copula-based joint deficit index for droughts. *Journal of Hydrology* 380(1-2):121-134.
- Kao SC and Govindaraju RS (2008) Trivariate statistical analysis of extreme rainfall events via the Plackett family of copulas. *Water Resources Research* 44(2).
- Karahacane H, Meddi M, Chebana F and Saaed HA (2020) Complete multivariate flood frequency analysis, applied to northern Algeria. *Journal of Flood Risk Management* 13(4):e12619.
- Karahan H, Ceylan H and Tamer Ayvaz M (2007) Predicting rainfall intensity using a genetic algorithm approach. *Hydrological Processes: An International Journal* 21(4):470-475.
- Kelly K and Krzysztofowicz R (1997) A bivariate meta-Gaussian density for use in hydrology. *Stochastic Hydrology and hydraulics* 11(1):17-31.
- Khaliq M, Ouarda T, Ondo J-C, Gachon P and Bobée B (2006) Frequency analysis of a sequence of dependent and/or non-stationary hydro-meteorological observations: A review. *Journal of hydrology* 329(3):534-552.
- Khan F, Spöck G and Pilz J (2019) A novel approach for modelling pattern and spatial dependence structures between climate variables by combining mixture models with copula models. *International Journal of Climatology*.
- Khedun CP, Mishra AK, Singh VP and Giardino JR (2014) A copula-based precipitation forecasting model: Investigating the interdecadal modulation of ENSO's impacts on monthly precipitation. *Water Resources Research* 50(1):580-600.
- Kim D, Kim J-M, Liao S-M and Jung Y-S (2013) Mixture of D-vine copulas for modeling dependence. *Computational Statistics & Data Analysis* 64:1-19.
- Kim G, Silvapulle MJ and Silvapulle P (2007) Comparison of semiparametric and parametric methods for estimating copulas. *Computational Statistics & Data Analysis* 51(6):2836-2850.
- Kjeldsen T and Prosdocimi I (2015a) A bivariate extension of the Hosking and Wallis goodness-of-fit measure for regional distributions. *Water Resources Research* 51(2):896-907.
- Kjeldsen T and Prosdocimi I (2015b) A bivariate extension of the Hosking and Wallis goodness-of-fit measure for regional distributions. *Water Resources Research* 51(2):896-907.
- Klein B, Schumann AH and Pahlow M (2011) Copulas—new risk assessment methodology for dam safety. *Flood Risk Assessment and Management*, Springer. p 149-185.
- Kojadinovic I, Quessy J-F and Rohmer T (2016) Testing the constancy of Spearman's rho in multivariate time series. *Annals of the Institute of Statistical Mathematics* 68(5):929-954.
- Kojadinovic I and Yan J (2010) Comparison of three semiparametric methods for estimating dependence parameters in copula models. *Insurance: Mathematics and Economics* 47(1):52-63.
- Kojadinovic I and Yan J (2011) A goodness-of-fit test for multivariate multiparameter copulas based on multiplier central limit theorems. *Statistics and Computing* 21(1):17-30.
- Kojadinovic I, Yan J and Holmes M (2011) Fast large-sample goodness-of-fit tests for copulas. *Statistica Sinica* 21(2):841.
- Kosmidis I and Karlis D (2016) Model-based clustering using copulas with applications. *Statistics and computing* 26(5):1079-1099.
- Kundzewicz ZW and Robson AJ (2004) Change detection in hydrological records—a review of the methodology/revue méthodologique de la détection de changements dans les chroniques hydrologiques. *Hydrological sciences journal* 49(1):7-19.

- Kysely J and Picek J (2007) Regional growth curves and improved design value estimates of extreme precipitation events in the Czech Republic. *Climate research* 33(3):243-255.
- Laio F (2004) Cramer–von Mises and Anderson-Darling goodness of fit tests for extreme value distributions with unknown parameters. *Water Resources Research* 40(9).
- Laio F, Di Baldassarre G and Montanari A (2009) Model selection techniques for the frequency analysis of hydrological extremes. *Water Resources Research* 45(7).
- Lall U (1995) Recent advances in nonparametric function estimation: Hydrologic applications. *Reviews of Geophysics* 33(S2):1093-1102.
- Lee T and Salas JD (2011) Copula-based stochastic simulation of hydrological data applied to Nile River flows. *Hydrology Research* 42(4):318-330.
- Leroux BG (1992) Consistent estimation of a mixing distribution. *The Annals of Statistics* :1350-1360.
- Leytham K (1984) Maximum likelihood estimates for the parameters of mixture distributions. *Water resources research* 20(7):896-902.
- Li C, Singh VP and Mishra AK (2012) Simulation of the entire range of daily precipitation using a hybrid probability distribution. *Water resources research* 48(3).
- Li C, Singh VP and Mishra AK (2013) A bivariate mixed distribution with a heavy-tailed component and its application to single-site daily rainfall simulation. *Water Resources Research* 49(2):767-789.
- Li J and Liu RY (2004) New nonparametric tests of multivariate locations and scales using data depth. *Statistical Science* :686-696.
- Liu G, Fu Y, Zhang J, Pu X and Wang B (2019) EM-test for homogeneity in a two-sample problem with a mixture structure. *Journal of Applied Statistics* :1-15.
- Liu RY and Singh K (1993) A quality index based on data depth and multivariate rank tests. *Journal of the American Statistical Association* 88(421):252-260.
- Lund R, Wang XL, Lu QQ, Reeves J, Gallagher C and Feng Y (2007) Changepoint detection in periodic and autocorrelated time series. *Journal of Climate* 20(20):5178-5190.
- Lung-Yut-Fong A, Lévy-Leduc C and Cappé O (2011) Homogeneity and change-point detection tests for multivariate data using rank statistics. *arXiv preprint arXiv:1107.1971*.
- Lung-Yut-Fong A, Lévy-Leduc C and Cappé O (2015) Homogeneity and change-point detection tests for multivariate data using rank statistics. *Journal de la Société Française de Statistique* 156(4):133-162.
- Ma M, Song S, Ren L, Jiang S and Song J (2013) Multivariate drought characteristics using trivariate Gaussian and Student t copulas. *Hydrological processes* 27(8):1175-1190.
- Mazouz R, Assani AA, Quessy J-F and Légaré G (2012) Comparison of the interannual variability of spring heavy floods characteristics of tributaries of the St. Lawrence River in Quebec (Canada). *Advances in water resources* 35:110-120.
- McLachlan G and Jones P (1988) Fitting mixture models to grouped and truncated data via the EM algorithm. *Biometrics* :571-578.
- Mesfioui M, Quessy JF and Toupin MH (2009) On a new goodness-of-fit process for families of copulas. *Canadian Journal of Statistics* 37(1):80-101.
- Meylan P, Favre A-C and Musy A (2012) *Predictive hydrology: a frequency analysis approach*. CRC Press,
- Milly PC, Betancourt J, Falkenmark M, Hirsch RM, Kundzewicz ZW, Lettenmaier DP, Stouffer RJ, Dettinger MD and Krysanova V (2015) On critiques of “Stationarity is dead: whither water management?”. *Water Resources Research* 51(9):7785-7789.

- Müller T, Schütze M and Bárdossy A (2017) Temporal asymmetry in precipitation time series and its influence on flow simulations in combined sewer systems. *Advances in Water Resources* 107:56-64.
- Nayak MA and Villarini G (2016) Evaluation of the capability of the Lombard test in detecting abrupt changes in variance. *Journal of Hydrology* 534:451-465.
- Nelsen RB (2006) *An Introduction to Copulas*. Springer Science & Business Media,
- Nelsen RB (2013) *An introduction to copulas*. Springer Science & Business Media,
- Nguyen C, Bhatti MI, Komorníková M and Komorník J (2016) Gold price and stock markets nexus under mixed-copulas. *Economic Modelling* 58:283-292.
- Oja H and Randles RH (2004) Multivariate nonparametric tests. *Statistical Science* :598-605.
- Ouarda T, Charron C, Kumar KN, Marpu P, Ghedira H, Molini A and Khayal I (2014) Evolution of the rainfall regime in the United Arab Emirates. *Journal of Hydrology* 514:258-270.
- Ouarda T and El-Adlouni S (2011) Bayesian nonstationary frequency analysis of hydrological variables. *JAWRA Journal of the American Water Resources Association* 47(3):496-505.
- Peterson TC, Easterling DR, Karl TR, Groisman P, Nicholls N, Plummer N, Torok S, Auer I, Boehm R and Gullett D (1998) Homogeneity adjustments of in situ atmospheric climate data: a review. *International journal of climatology* 18(13):1493-1517.
- Pettitt A (1979) A non-parametric approach to the change-point problem. *Applied statistics* :126-135.
- Pickands J (1989) Multivariate negative exponential and extreme value distributions. *Extreme Value Theory*, Springer. p 262-274.
- Poulin A, Huard D, Favre A-C and Pugin S (2007) Importance of tail dependence in bivariate frequency analysis. *Journal of Hydrologic Engineering* 12(4):394-403.
- Qu L and Lu Y (2019) Copula density estimation by finite mixture of parametric copula densities. *Communications in Statistics-Simulation and Computation* :1-23.
- Quessy J-F (2019) Consistent nonparametric tests for detecting gradual changes in the marginals and the copula of multivariate time series. *Statistical Papers* 60(3):367-396.
- Quessy JF, Saïd M and Favre AC (2013) Multivariate Kendall's tau for change-point detection in copulas. *Canadian Journal of Statistics* 41(1):65-82.
- Rao AR and Hamed KH (2000) Flood frequency analysis. *Ingeniería del Agua* 7(3):309-309.
- Reca J, Martínez J, Gil C and Baños R (2008) Application of several meta-heuristic techniques to the optimization of real looped water distribution networks. *Water Resources Management* 22(10):1367-1379.
- Reddy MJ and Singh VP (2014) Multivariate modeling of droughts using copulas and meta-heuristic methods. *Stochastic environmental research and risk assessment* 28(3):475-489.
- Reeves J, Chen J, Wang XL, Lund R and Lu QQ (2007) A review and comparison of changepoint detection techniques for climate data. *Journal of Applied Meteorology and Climatology* 46(6):900-915.
- Requena A, Mediero L and Garrote L (2013a) A bivariate return period based on copulas for hydrologic dam design: accounting for reservoir routing in risk estimation. *Hydrology and Earth system sciences* 17(8):3023.
- Requena A, Mediero Orduña L and Garrote de Marcos L (2013b) A bivariate return period based on copulas for hydrologic dam design: accounting for reservoir routing in risk estimation. *Hydrology and Earth System Sciences* 17(8):3023-3038.
- Requena AI, Chebana F and Mediero L (2016) A complete procedure for multivariate index-flood model application. *Journal of Hydrology* 535:559-580.

- Ribes A, Zwiers FW, Azaïs J-M and Naveau P (2017) A new statistical approach to climate change detection and attribution. *Climate Dynamics* 48(1-2):367-386.
- Rougé C, Ge Y and Cai X (2013) Detecting gradual and abrupt changes in hydrological records. *Advances in Water Resources* 53:33-44.
- Sadegh M, Ragno E and AghaKouchak A (2017) Multivariate Copula Analysis Toolbox (MvCAT): Describing dependence and underlying uncertainty using a Bayesian framework. *Water Resources Research*.
- Sadegh M, Vrugt JA, Xu C and Volpi E (2015) The stationarity paradigm revisited: Hypothesis testing using diagnostics, summary metrics, and DREAM (ABC). *Water Resources Research* 51(11):9207-9231.
- Salvadori G and De Michele C (2004) Frequency analysis via copulas: Theoretical aspects and applications to hydrological events. *Water Resources Research* 40(12).
- Salvadori G and De Michele C (2010) Multivariate multiparameter extreme value models and return periods: A copula approach. *Water resources research* 46(10).
- Salvadori G, De Michele C, Kottegoda NT and Rosso R (2007) *Extremes in nature: an approach using copulas*. Springer Science & Business Media,
- Salvadori G, Durante F, De Michele C and Bernardi M (2018) Hazard assessment under multivariate distributional change-points: Guidelines and a flood case study. *Water* 10(6):751.
- Salvadori G, Durante F, De Michele C, Bernardi M and Petrella L (2016) A multivariate copula-based framework for dealing with hazard scenarios and failure probabilities. *Water Resources Research* 52(5):3701-3721.
- Salvadori G and Michele CD (2011) Estimating strategies for multiparameter multivariate extreme value copulas. *Hydrology and Earth System Sciences* 15(1):141-150.
- Santhosh D and Srinivas V (2013) Bivariate frequency analysis of floods using a diffusion based kernel density estimator. *Water Resources Research* 49(12):8328-8343.
- Sarhadi A, Burn DH, Concepción Ausín M and Wiper MP (2016) Time varying nonstationary multivariate risk analysis using a dynamic Bayesian copula. *Water Resources Research*.
- Schoelzel C and Friederichs P (2008) Multivariate non-normally distributed random variables in climate research—introduction to the copula approach. *Nonlin. Processes Geophys.* 15(5):761-772.
- Schumann AH (2011) *Flood Risk Assessment and Management: How to Specify Hydrological Loads, Their Consequences and Uncertainties*. Springer Science & Business Media,
- Seidou O, Asselin J and Ouarda T (2007) Bayesian multivariate linear regression with application to change point models in hydrometeorological variables. *Water Resources Research* 43(8).
- Seidou O and Ouarda TB (2007) Recursion-based multiple changepoint detection in multiple linear regression and application to river streamflows. *Water Resources Research* 43(7).
- Serfling R and Xiao P (2007) A contribution to multivariate L-moments: L-comoment matrices. *Journal of Multivariate Analysis* 98(9):1765-1781.
- Serinaldi F, Kilsby CG and Lombardo F (2018) Untenable nonstationarity: An assessment of the fitness for purpose of trend tests in hydrology. *Advances in Water Resources* 111:132-155.
- Shiau J-T, Wang H-Y and Tsai C-T (2006) Bivariate frequency analysis of floods using copulas. *Journal of the American Water Resources Association* 42(6):1549-1564.
- Shih JH (1998) A goodness-of-fit test for association in a bivariate survival model. *Biometrika* 85(1):189-200.
- Shih JH and Louis TA (1995) Inferences on the association parameter in copula models for bivariate survival data. *Biometrics* :1384-1399.

- Shin J-Y, Heo J-H, Jeong C and Lee T (2014) Meta-heuristic maximum likelihood parameter estimation of the mixture normal distribution for hydro-meteorological variables. *Stochastic environmental research and risk assessment* 28(2):347-358.
- Shin J-Y, Lee T and Ouarda TB (2015) Heterogeneous mixture distributions for modeling multisource extreme rainfalls. *Journal of Hydrometeorology* 16(6):2639-2657.
- Shin J-Y, Ouarda TB and Lee T (2016) Heterogeneous mixture distributions for modeling wind speed, application to the UAE. *Renewable Energy* 91:40-52.
- Sinclair C, Spurr B and Ahmad M (1990) Modified anderson darling test. *Communications in Statistics-Theory and Methods* 19(10):3677-3686.
- Singh V, Wang S and Zhang L (2005a) Frequency analysis of nonidentically distributed hydrologic flood data. *Journal of Hydrology* 307(1-4):175-195.
- Singh V, Wang S and Zhang L (2005b) Frequency analysis of nonidentically distributed hydrologic flood data. *Journal of Hydrology* 307(1):175-195.
- Singh VP and Zhang L (2018) Copula–entropy theory for multivariate stochastic modeling in water engineering. *Geoscience Letters* 5(1):6.
- Sivanandam S and Deepa S (2008) Genetic algorithms. *Introduction to genetic algorithms*, Springer. p 15-37.
- Sklar A (1959) Fonctions de répartition à n dimensions et leurs marges. *Publications se l'Institut de Statistique de Paris* 8:229-231.
- Smith JA, Villarini G and Baeck ML (2011) Mixture distributions and the hydroclimatology of extreme rainfall and flooding in the eastern United States. *Journal of Hydrometeorology* 12(2):294-309.
- Song S and Singh VP (2010) Frequency analysis of droughts using the Plackett copula and parameter estimation by genetic algorithm. *Stochastic Environmental Research and Risk Assessment* 24(5):783-805.
- Taskinen S, Oja H and Randles RH (2005) Multivariate nonparametric tests of independence. *Journal of the American Statistical Association* 100(471):916-925.
- Thongkairat S, Yamaka W and Sriboonchitta S (2019) Bayesian Approach for Mixture Copula Model. *Beyond Traditional Probabilistic Methods in Economics*. (Cham, 2019//), Kreinovich V, Thach NN, Trung ND and Van Thanh D (Édit.) Springer International Publishing, p 818-827.
- Van der Vaart A (1996) Efficient maximum likelihood estimation in semiparametric mixture models. *The Annals of Statistics* 24(2):862-878.
- Vandenberghe S, Verhoest N and De Baets B (2010) Fitting bivariate copulas to the dependence structure between storm characteristics: A detailed analysis based on 105 year 10 min rainfall. *Water resources research* 46(1).
- Vezzoli R, Salvadori G and De Michele C (2017) A distributional multivariate approach for assessing performance of climate-hydrology models. *Scientific reports* 7(1):12071.
- Villarini G and Smith JA (2010) Flood peak distributions for the eastern United States. *Water Resources Research* 46(6).
- Villarini G, Smith JA, Serinaldi F, Bales J, Bates PD and Krajewski WF (2009) Flood frequency analysis for nonstationary annual peak records in an urban drainage basin. *Advances in Water Resources* 32(8):1255-1266.
- Vittal H, Singh J, Kumar P and Karmakar S (2015) A framework for multivariate data-based at-site flood frequency analysis: Essentiality of the conjugal application of parametric and nonparametric approaches. *Journal of Hydrology* 525:658-675.

- Vogel D and Fried R (2015) Robust change detection in the dependence structure of multivariate time series. *Modern Nonparametric, Robust and Multivariate Methods*, Springer. p 265-288.
- Volpi E and Fiori A (2014) Hydraulic structures subject to bivariate hydrological loads: Return period, design, and risk assessment. *Water Resources Research* 50(2):885-897.
- Von Storch H and Zwiers FW (2001) *Statistical analysis in climate research*. Cambridge university press,
- Vrac M, Billard L, Diday E and Chédin A (2012) Copula analysis of mixture models. *Computational Statistics* 27(3):427-457.
- Vrac M, Chédin A and Diday E (2005) Clustering a global field of atmospheric profiles by mixture decomposition of copulas. *Journal of Atmospheric and Oceanic Technology* 22(10):1445-1459.
- Vrac M and Naveau P (2007) Stochastic downscaling of precipitation: From dry events to heavy rainfalls. *Water resources research* 43(7).
- Wang W and Wells MT (2000) Model selection and semiparametric inference for bivariate failure-time data. *Journal of the American Statistical Association* 95(449):62-72.
- White H (1981) Consequences and detection of misspecified nonlinear regression models. *Journal of the American Statistical Association* 76(374):419-433.
- Wied D, Dehling H, Van Kampen M and Vogel D (2014) A fluctuation test for constant Spearman's rho with nuisance-free limit distribution. *Computational Statistics & Data Analysis* 76:723-736.
- Wilks DS (2011) *Statistical methods in the atmospheric sciences*. Academic press,
- Willems P (2013) Revision of urban drainage design rules after assessment of climate change impacts on precipitation extremes at Uccle, Belgium. *Journal of Hydrology* 496:166-177.
- Woodward WA, Parr WC, Schucany WR and Lindsey H (1984) A comparison of minimum distance and maximum likelihood estimation of a mixture proportion. *Journal of the American Statistical Association* 79(387):590-598.
- Xie H, Li D and Xiong L (2014) Exploring the ability of the Pettitt method for detecting change point by Monte Carlo simulation. *Stochastic environmental research and risk assessment* 28(7):1643-1655.
- Xiong L, Jiang C, Xu CY, Yu Kx and Guo S (2015) A framework of change-point detection for multivariate hydrological series. *Water Resources Research* 51(10):8198-8217.
- Yan L, Xiong L, Liu D, Hu T and Xu CY (2016) Frequency analysis of nonstationary annual maximum flood series using the time-varying two-component mixture distributions. *Hydrological Processes*.
- Yan L, Xiong L, Liu D, Hu T and Xu CY (2017) Frequency analysis of nonstationary annual maximum flood series using the time-varying two-component mixture distributions. *Hydrological Processes* 31(1):69-89.
- Yan L, Xiong L, Ruan G, Xu C-Y, Yan P and Liu P (2019) Reducing uncertainty of design floods of two-component mixture distributions by utilizing flood timescale to classify flood types in seasonally snow covered region. *Journal of Hydrology* 574:588-608.
- Yu J, Chen K, Mori J and Rashid MM (2013) A Gaussian mixture copula model based localized Gaussian process regression approach for long-term wind speed prediction. *Energy* 61:673-686.
- Yue S, Ouarda T, Bobée B, Legendre P and Bruneau P (1999a) The Gumbel mixed model for flood frequency analysis. *Journal of hydrology* 226(1-2):88-100.

- Yue S, Ouarda T, Bobée B, Legendre P and Bruneau P (1999b) The Gumbel mixed model for flood frequency analysis. *Journal of hydrology* 226(1):88-100.
- Yue S, Ouarda TBMJ and Bobée B (2001) A review of bivariate gamma distributions for hydrological application. *Journal of Hydrology* 246(1–4):1-18.
- Yue S, Pilon P and Cavadias G (2002) Power of the Mann–Kendall and Spearman's rho tests for detecting monotonic trends in hydrological series. *Journal of hydrology* 259(1):254-271.
- Yuille AL, Stolorz P and Utans J (1994) Statistical physics, mixtures of distributions, and the EM algorithm. *Neural Computation* 6(2):334-340.
- Zardasht V, Parsi S and Mousazadeh M (2015) On empirical cumulative residual entropy and a goodness-of-fit test for exponentiality. *Statistical Papers* 56(3):677-688.
- Zhang L and Singh V (2006) Bivariate flood frequency analysis using the copula method. *Journal of Hydrologic Engineering*.
- Zhang L and Singh VP (2007a) Bivariate rainfall frequency distributions using Archimedean copulas. *Journal of Hydrology* 332(1-2):93-109.
- Zhang L and Singh VP (2007b) Trivariate flood frequency analysis using the Gumbel–Hougaard copula. *Journal of Hydrologic Engineering* 12(4):431-439.
- Zhang L and Singh VP (2012) Bivariate rainfall and runoff analysis using entropy and copula theories. *Entropy* 14(9):1784-1812.
- Zhang L and Singh VP (2019) *Copulas and their applications in water resources engineering*. Cambridge University Press,
- Zheng X and Katz RW (2008) Mixture model of generalized chain-dependent processes and its application to simulation of interannual variability of daily rainfall. *Journal of Hydrology* 349(1-2):191-199.