

# SIMULATION OF MULTIVARIATE EXTREMES: A WASSERSTEIN- AITCHISON GAN APPROACH

Stéphane Lhaut, Holger Rootzén, Johan  
Segers

REPRINT | 2026 / 07

## **ISBA**

Voie du Roman Pays 20 - L1.04.01

B-1348 Louvain-la-Neuve

Email : [lidam-library@uclouvain.be](mailto:lidam-library@uclouvain.be)

<https://uclouvain.be/en/research-institutes/lidam/isba/publication>

# Simulation of Multivariate Extremes: a Wasserstein–Aitchison GAN approach

Stéphane Lhaut  <sup>\*1,4</sup>, Holger Rootzén  <sup>†2</sup>, and Johan Segers  <sup>‡1,3</sup>

<sup>1</sup>Institute of Statistics, Biostatistics and Actuarial Sciences, LIDAM, UCLouvain, Voie du Roman Pays, 20, 1348 Louvain-La-Neuve, Belgium

<sup>2</sup>Department of Mathematical Sciences, Chalmers University of Technology and University of Gothenburg, Gothenburg, Sweden

<sup>3</sup>Department of Mathematics, KU Leuven, Celestijnenlaan 200B, 3001 Heverlee, Belgium

<sup>4</sup>CREST, ENSAE Paris, 5 avenue Henry Le Chatelier 91120 Palaiseau, France

January 13, 2026

## Abstract

Economically responsible mitigation of multivariate extreme risks—such as extreme rainfall over large areas, large simultaneous variations in many stock prices, or widespread breakdowns in transportation systems—requires assessing the resilience of the systems under plausible stress scenarios. This paper uses Extreme Value Theory (EVT) to develop a new approach to simulating such multivariate extreme events. Specifically, we assume that after transformation to a standard scale the distribution of the random phenomenon of interest is multivariate regular varying and use this to provide a sampling procedure for extremes on the original scale. Our procedure combines a Wasserstein–Aitchison Generative Adversarial Network (WA-GAN) to simulate the tail dependence structure on the standard scale with joint modeling of the univariate marginal tails on the original scale. The WA-GAN procedure relies on the angular measure—encoding the distribution on the unit simplex of the angles of extreme observations—after transformation to Aitchison coordinates, which allows the Wasserstein-GAN algorithm to be run in a linear space. Our method is applied both to simulated data under various tail dependence scenarios and to a financial data set from the Kenneth French Data Library. The proposed algorithm demonstrates strong performance compared to existing alternatives in the literature, both in capturing tail dependence structures and in generating accurate new extreme observations.

**Keywords:** Aitchison coordinates, Angular measure, Extreme value theory, Generative adversarial networks, Generative AI for extremes, Multivariate analysis, Wasserstein distance

---

<sup>\*</sup>stephane.lhaut@ensae.fr, corresponding author

<sup>†</sup>hrootzen@chalmers.se

<sup>‡</sup>jjjsegers@kuleuven.be

# 1 Introduction

**Multivariate EVT.** Dependence between extreme events is often of crucial importance. Extreme rainfall is even more dangerous if it occurs simultaneously at many spatial locations; costs associated with breakdowns in transportation systems are more harmful if they occur in many parts of the system; and extreme fluctuations in financial markets can cause much more damage if they occur across many sectors of the economy. Over the last decades, statistical research on dependent, multivariate extremes has grown rapidly and now includes many useful, mostly parametric, models and methods. Often, the computational challenges associated with estimating the parameters of these models restrict their use to relatively low-dimensional settings. When spatial information can be incorporated into the model, dimensions may be higher, but this is limited to specific applications, mainly climatic ones. For background on EVT, we refer to standard monographs such as [7, 15, 48] or more recent ones such as [39, 49].

**GenAI.** Generative Artificial Intelligence (GenAI) methods based on artificial neural networks now provide completely new opportunities to address challenges in high dimensions. The GenAI approach differs from the standard statistical modeling perspective in that the output of these methods is a sampling algorithm rather than a set of parameters from a given parametric model, from which sampling can be performed in a second stage. Here, no explicit assumption on the data distribution is made, even though certain assumptions may help in designing the neural network architectures. One of the most popular such algorithms is the Generative Adversarial Network (GAN) [26] or its variant, the Wasserstein-GAN (WGAN) [5]. See [10, Chapter 17] for a pedagogical introduction to GANs and their variations. Basically, the idea is to train two neural networks simultaneously: one aims to generate new realizations of the random phenomenon of interest (the generator) by transforming samples from a known (latent) distribution, while the other attempts to discriminate (the discriminator) between the true and generated realizations. The difference between GAN and WGAN lies in the distance used to compare samples, which is related to the Jensen–Shannon divergence in the case of GAN and to the Wasserstein distance in the case of WGAN. The latter aims to avoid the well-known “mode collapse” issues associated with GANs and is therefore often used as the default choice in such approaches [5]. Initially, GAN algorithms were proposed for synthetic image generation, but they have since been applied to many types of data, in particular tabular data, which are of special interest to statisticians [4]. Some theoretical properties of the WGAN algorithm can be found in [9]. Many other sampling algorithms exist in the literature, such as Normalizing Flows (NF), Variational Autoencoders (VAE), or Diffusion Models (DM), each with its own merits and drawbacks. As our method is based on the WGAN algorithm, we do not discuss these alternative approaches in detail here, but the interested reader may find an introduction to them in [10, Chapters 18, 19, 20].

**Bridging EVT and GenAI.** Using GenAI approaches to simulate new multivariate extreme scenarios is challenging for multiple reasons. First, it is well known from EVT



that the tail dependence structure of a random vector may differ drastically from its central dependence. As such, a straightforward application of a sampling algorithm to the entire dataset has little chance of accurately capturing the tail behavior of the underlying data-generating process. Second, as pointed out in [25], models based on transformations of light-tailed noise, such as GANs, fail to capture the tail behavior of heavy-tailed distributions, and these algorithms must therefore be adapted. Lastly, GenAI approaches typically require large amounts of data, whereas extreme events are, by definition, rare. Consequently, the algorithms need to be adapted in order to deal with the limited amount of data relevant for tail problems.

The use of EVT methods within GenAI frameworks for simulating accurate new extremes has attracted considerable attention recently. Adaptations of GAN algorithms have been considered in [2, 3, 8, 33]. These approaches all rely on univariate EVT methods applied to each margin, but they do not incorporate tail dependence structures arising from multivariate EVT. An inherently multivariate approach is proposed in [11], which applies univariate EVT to the margins and models the dependence of the full dataset after a copula transformation. No multivariate EVT is used there, as is also the case in [25], whose Heavy-Tailed GAN (HTGAN) method aims to model dependence after transformation to heavy-tailed margins using a heavy-tailed latent distribution for the generator. GAN approaches that explicitly make use of multivariate EVT can be found in the Generalized Pareto GAN (GPGAN) algorithm [37], which targets multivariate threshold exceedances based on multivariate Generalized Pareto (MGP) distributions [52], or in the  $d$ -max-decreasing Neural Networks (dMNN) approach [28], which relies on multivariate block maxima and the Pickands dependence function, see [7, Chapter 8] or [15, Chapter 6].

Other approaches to sampling multivariate extremes can be found in the literature, some of which are very recent. A VAE-based method is proposed in [36]. Normalizing flow approaches in the context of “geometric extremes” (a framework different from multivariate regular variation) are presented in [32] and [41]. Issues related to generating heavy-tailed outputs using normalizing flows are also discussed in [31]. We will not compare our method to these approaches, as our focus is on GAN algorithms and their variants. A broader comparison of many existing generative methods for extremes is provided in [57].

**WA-GAN.** The method we propose, Wasserstein–Aitchison GAN (WA-GAN), also relies on multivariate EVT and takes advantage of the well-known polar decomposition of a multivariate regularly varying random vector into an independent radial component and an angular component, which we represent by a point on the unit simplex. Since we work on a standard scale, no assumptions on the margins are required apart from continuity. The generative component of our method operates on the angular values. The advantage of this approach is that, at the angular scale, the data are bounded, and we can rely on Aitchison coordinates, a tool from compositional data analysis, to simulate values on the unit simplex by applying a WGAN on a linear space. The generative algorithm is then combined with standard marginal EVT modeling to yield a generative algorithm for multivariate threshold exceedances within the MGP framework.

We compare our approach to GPGAN, which follows a similar strategy but models both the marginal tails and the tail dependence structure simultaneously, and to HTGAN, which uses the same standardization as our method but then applies a standard GAN with heavy-tailed input to model dependence, without exploiting the specific structure of tail dependence.

**Outline.** Section 2 presents the background required to introduce our WA-GAN algorithm, together with a theoretical result (Proposition 1) that provides the tools needed to transform angular data on the simplex into new multivariate extreme scenarios. Section 3 describes the details of WA-GAN and presents the corresponding algorithms. Section 4.1 describes how performance of algorithms was evaluated and the architecture used for WA-GAN. In Section 4.2, WA-GAN is compared to HTGAN [25] and GPGAN [37] using simulated data with logistic and Hüsler–Reiss dependence structures [7, Section 9.2.2] in dimensions  $d \in \{10, 20, 50\}$ . Section 4.2.3 discusses how the use of non-Gaussian latent distributions for the generator can help improve results when asymptotic independence is suspected. In Section 4.3, we apply WA-GAN to financial data consisting of daily returns for  $d = 30$  industry portfolios. Section 5 summarizes the results.

## 2 Background

### 2.1 Regular variation and extreme value analysis

**Multivariate regular variation and angular measure.** Let  $\mathbf{X} = (X_1, \dots, X_d)$  be the random vector of interest with values in  $\mathbb{R}^d$ . Based on a random sample of the (unknown) distribution of  $\mathbf{X}$ , we wish to generate new samples with the same tail behavior as  $\mathbf{X}$ , even at levels beyond those encountered in the sample. Such extrapolation is possible only if we are willing to make some regularity assumptions. Our core assumptions to obtain accurate modeling of the tail of  $\mathbf{X}$  are founded in the theory of multivariate regular variation [47]. There are two approaches to use this hypothesis.

1. Either we postulate it on the random vector  $\mathbf{X}$  itself. This implies that all univariate tail indices are same.
2. Or we postulate it on a transformed vector  $\mathbf{V} = v(\mathbf{X})$ , where  $v$  is a coordinatewise transformation of  $\mathbf{X}$  so that the  $d$  components of  $\mathbf{V}$  all have the same distribution. This implies that the multivariate assumption only depends on the dependence structure of  $\mathbf{X}$  and not on the margins, which can be assessed separately.

Here we choose the second approach and transform to unit-Pareto margins using the transformation

$$v : \mathbf{x} = (x_j)_{j=1}^d \mapsto v(\mathbf{x}) \stackrel{\text{def}}{=} \left( \frac{1}{1 - F_j(x_j)} \right)_{j=1}^d,$$

where  $F_j$  denotes the distribution function of  $X_j$ . We assume throughout that these marginal distribution functions are continuous. It is straightforward to see that after this

transformation the margins of  $\mathbf{V} = v(\mathbf{X})$  have unit-Pareto distributions,  $\mathbf{P}(V_j > y) = 1/y$  for  $y \geq 1$ . Our first assumption describes the asymptotic dependence structure of  $\mathbf{V}$ .

**Assumption 1** (Multivariate regular variation). *There exists a non-zero Borel measure  $\nu$  on  $\mathbb{E} \stackrel{\text{def}}{=} [0, \infty)^d \setminus \{\mathbf{0}\}$  which is finite on Borel sets bounded away from  $\mathbf{0}$  and such that*

$$\lim_{t \rightarrow \infty} t \mathbf{P}(t^{-1} \mathbf{V} \in B) = \nu(B)$$

for all Borel sets  $B$  of  $\mathbb{E}$  bounded away from  $\mathbf{0}$  with  $\nu(\partial B) = 0$ , with  $\partial B$  the topological boundary of  $B$ .

Intuitively, the measure  $\nu$  helps in describing the tail behavior of  $\mathbf{V}$  as for large  $t > 0$ , we have  $\mathbf{P}(\mathbf{V} \in tB) \approx \nu(B)/t$  where  $tB \stackrel{\text{def}}{=} \{t\mathbf{x} : \mathbf{x} \in B\}$ . The measure  $\nu$  is often referred to as the *exponent measure* of  $\mathbf{V}$ , because of its appearance in the exponent of the expression of the multivariate extreme value distribution to which the law of  $\mathbf{V}$  is attracted [15, Definition 6.1.7].

Our procedure below will benefit from an equivalent formulation of Assumption 1 in polar coordinates with respect to the  $L_1$ -norm  $|\mathbf{x}|_1 \stackrel{\text{def}}{=} \sum_{j=1}^d |x_j|$  for  $\mathbf{x} \in \mathbb{R}^d$ . More concretely, it can be shown [48, Theorem 6.1] that due to a convenient homogeneity property of  $\nu$  [15, Theorem 6.1.9], if we put

$$R \stackrel{\text{def}}{=} |\mathbf{V}|_1 \quad \text{and} \quad \mathbf{W} \stackrel{\text{def}}{=} R^{-1} \mathbf{V}, \quad (1)$$

where  $\mathbf{W}$  takes values in the unit simplex  $\Delta_{d-1} \stackrel{\text{def}}{=} \{\mathbf{x} \in [0, 1]^d : |\mathbf{x}|_1 = 1\}$ , there exists a probability measure  $\Phi$  on  $\Delta_{d-1}$  such that

$$\lim_{t \rightarrow \infty} \mathbf{P}(\mathbf{W} \in A \mid R \geq t) = \Phi(A), \quad (2)$$

for any Borel set  $A \subseteq \Delta_{d-1}$  such that  $\Phi(\partial A) = 0$ .

The measure  $\Phi$  is often referred to as the *angular (probability) measure* of  $\mathbf{V}$  as it describes how its “angular” component  $\mathbf{W}$  behaves when its radial part  $R$  becomes large. The set of possible probability measures that can be obtained in such a way forms a nonparametric class as the only constraint they should satisfy—due to the unit-Pareto margins of  $\mathbf{V}$ —are

$$\int_{\Delta_{d-1}} w_j d\Phi(\mathbf{w}) = \frac{1}{d}, \quad j \in \{1, \dots, d\}, \quad (3)$$

which in turn implies

$$\mathbf{P}(R > t) \sim \frac{d}{t}, \quad \text{as } t \rightarrow \infty. \quad (4)$$

It is possible for  $\Phi$  to have positive mass on one of the faces of the simplex  $\Delta_{d-1}$ , that is, sets for which at least one of the coordinates is zero. In terms of extremes, it corresponds to scenarios where some components of  $\mathbf{X}$  are large while others are small. These are called *extreme directions* and have to be treated by specific methods which our methodology is not able to handle, see, e.g., [42]. We exclude this situation here.

**Assumption 2.** *The measure  $\Phi$  in (2) is concentrated on the open simplex  $\Delta_{d-1}^\circ \stackrel{\text{def}}{=} \{\mathbf{x} \in (0, 1)^d : |\mathbf{x}|_1 = 1\}$ , that is, if  $\Theta \sim \Phi$ , then*

$$\mathbb{P}(\Theta \in \Delta_{d-1}^\circ) = 1.$$

**Univariate generalized Pareto distributions.** Assumption 1 is not sufficient by itself to model the tail behavior of  $\mathbf{X}$  as it only concerns the dependence structure. Our second main assumption concerns the tails of the margins  $X_j$  and allows for heterogeneous tails. Here and below,  $u_{*j}$  denotes the (possibly infinite) upper endpoint of  $X_j$ .

**Assumption 3** (Generalized Pareto limit of marginal tails). *For every  $j \in \{1, \dots, d\}$ , there exist  $\xi_j \in \mathbb{R}$  and a function  $\sigma_j(\cdot) : (-\infty, u_{*j}) \rightarrow (0, \infty)$  such that for all  $y > 0$ ,*

$$\lim_{u \nearrow u_{*j}} \mathbb{P}\left(\frac{X_j - u}{\sigma_j(u)} > y \mid X_j > u\right) = (1 + \xi_j y)_+^{-1/\xi_j},$$

where  $a_+ \stackrel{\text{def}}{=} \max(a, 0)$  for  $a \in \mathbb{R}$ ; if  $\xi_j = 0$ , the limit is to be understood as  $\exp(-y)$ .

By the Pickands–Balkema–de Haan theorem [6, 45], Assumption 3 is equivalent to the assumption that  $X_j$  is in the maximum domain of attraction of a GEV distribution. The limit in Assumption 3 holds uniformly in  $y > 0$ . The assumption justifies the following approximation for  $y > 0$  such that  $1 + \xi_j y / \sigma_j > 0$ :

$$\mathbb{P}(X_j - u > y \mid X_j > u) \approx \left(1 + \frac{\xi_j y}{\sigma_j}\right)_+^{-1/\xi_j}, \quad (5)$$

where  $\sigma_j > 0$  is a scale parameter that accounts for the unknown function  $\sigma_j(\cdot)$ . Estimation of the bivariate parameter  $(\sigma_j, \xi_j)$  is typically performed by the method of maximum likelihood. This is the well-known *Peaks-Over-Threshold (PoT)* approach for modeling the tails of real-valued data, see for instance [54] and [7, Section 5.3]. A random variable with distribution function

$$H_{\sigma, \xi}(y) \stackrel{\text{def}}{=} 1 - \left(1 + \frac{\xi y}{\sigma}\right)_+^{-1/\xi}, \quad y > 0,$$

for some  $(\sigma, \xi) \in (0, \infty) \times \mathbb{R}$  is said to have a *generalized Pareto (GP) distribution*.

high-risk scenarios is the case where

*Remark 1* (Equivalent formulations of Assumption 3). As already explained, Assumption 3 is equivalent to the hypothesis that  $X_j$  is in the maximum domain of attraction of a GEV distribution. As a consequence,  $X_j$  satisfies all the equivalent formulations of this condition as presented, e.g., in Theorem 1.1.6 of [15]. As some of them will be needed below, we recall them here. Let  $b_j(t) \stackrel{\text{def}}{=} F_j^{-1}(1 - 1/t) = (1/(1 - F_j))^{-1}(t)$ , for  $t > 1$ , be the tail quantile function of  $F_j$ . Then, the following statements are equivalent to Assumption 3.

(i) There exist  $\xi_j \in \mathbb{R}$  and a scaling function  $a_j(\cdot) : (1, \infty) \rightarrow (0, \infty)$  such that

$$\lim_{t \rightarrow \infty} \mathbb{P} \left( \frac{X_j - b_j(t)}{a_j(t)} > y \mid X_j > b_j(t) \right) = (1 + \xi_j y)_+^{-1/\xi_j}, \quad y > 0.$$

(ii) There exist  $\xi_j \in \mathbb{R}$  and a scaling function  $a_j(\cdot) : (1, \infty) \rightarrow (0, \infty)$  such that

$$\lim_{t \rightarrow \infty} \frac{b_j(ty) - b_j(t)}{a_j(t)} = \frac{y^{\xi_j} - 1}{\xi_j}, \quad y > 0,$$

where the right-hand side is to be interpreted as  $\log y$  if  $\xi_j = 0$ . The convergence also holds locally uniformly in  $y \in (0, \infty)$ , see Section 1.2 in [15].

The scalar  $\xi_j \in \mathbb{R}$  in (i) and (ii) is equal to the one in Assumption 3, while the scaling function  $a_j(\cdot)$  in (i) and (ii) is related to  $\sigma_j(\cdot)$  in Assumption 3 via  $\sigma_j(u) = a_j(1/(1-F_j(u)))$  for  $u < u_{*j}$ .

**Multivariate generalized Pareto distributions.** What do Assumptions 1–3 tell us about the tail behavior of  $\mathbf{X}$ ? The answer can be conveniently formulated in terms of multivariate generalized Pareto (MGP) distributions, which extend Assumption 3 to the multivariate case. These distributions have been introduced in [52] and are reviewed in [43]. Below, we provide an introduction to their usage in modeling multivariate extremes along with a theoretical result (Proposition 1) which will be useful for sampling from a MGP distribution.

In arbitrary dimension, it is not clear how to define an “extreme” point. Some authors are interested in modeling the vector of componentwise maxima of the data, leading to the multivariate GEV distributions, studied extensively in the literature after the pioneering work of [16]. Recently, see e.g. [51], more attention has been given to an extension of the standard PoT procedure in a multivariate context. The approach relies on the fact that the excess of  $\mathbf{X}$  above a large threshold vector  $\mathbf{u}$  asymptotically follows an MGP distribution, with an appropriate definition of when an exceedance over a multivariate threshold takes place.

More concretely, given a vector  $\mathbf{u} \in \mathbb{R}^d$  of “high thresholds”, i.e., below but close to  $\mathbf{u}_* = (u_{*j})_{j=1}^d$ , where  $u_{*j}$  denotes the (possibly infinite) upper endpoint of  $X_j$ , we consider the random vector of excesses

$$\mathbf{Y}_{\mathbf{u}} \stackrel{\text{def}}{=} \mathbf{X} - \mathbf{u} \mid \mathbf{X} \not\leq \mathbf{u}, \quad (6)$$

where  $\mathbf{X} - \mathbf{u} \stackrel{\text{def}}{=} (X_j - u_j)_{j=1}^d$  and where  $\not\leq$  means that, for at least one  $j \in \{1, \dots, d\}$ , we have  $X_j > u_j$ . Here and below, operations on real numbers are extended to vectors in  $\mathbb{R}^d$  in a coordinate-wise manner. Under the aforementioned assumptions, we can establish the asymptotic distribution of  $\mathbf{Y}_{\mathbf{u}}$  as  $\mathbf{u} \rightarrow \mathbf{u}_*$ , in a particular way made precise in the statement of the result. Weak convergence in  $\mathbb{R}^d$  is denoted by  $\xrightarrow{\mathcal{L}}$ .

**Proposition 1** (Weak convergence of excesses to the MGP distribution). *Under Assumptions 1–3, we have*

$$\frac{\mathbf{X} - \mathbf{b}(t)}{\mathbf{a}(t)} \Big| \mathbf{X} \not\leq \mathbf{b}(t) \xrightarrow{\mathcal{L}} \frac{\mathbf{Y}^\xi - 1}{\xi}, \quad t \rightarrow \infty, \quad (7)$$

where  $\mathbf{Y}$  is a multivariate Pareto random vector, with distribution

$$\mathbf{Y} \stackrel{\mathcal{L}}{=} (Y\boldsymbol{\Theta} \mid Y\boldsymbol{\Theta} \not\leq 1) \quad (8)$$

with  $Y$  a unit-Pareto random variable independent of  $\boldsymbol{\Theta} \sim \Phi$ , and where  $\boldsymbol{\xi} = (\xi_j)_{j=1}^d$  is the vector of marginal coefficients as in Assumption 3;  $\mathbf{a}(\cdot) = (a_j(\cdot))_{j=1}^d$  and  $\mathbf{b}(\cdot) = (b_j(\cdot))_{j=1}^d$  are the same functions as the one introduced in Remark 1. Consequently, along the curve  $\mathbf{u} : t \in (1, \infty) \mapsto \mathbf{u}(t) \stackrel{\text{def}}{=} \mathbf{b}(t)$  we have,

$$\frac{\mathbf{Y}_{\mathbf{u}(t)}}{\boldsymbol{\sigma}(\mathbf{u}(t))} \xrightarrow{\mathcal{L}} \frac{\mathbf{Y}^\xi - 1}{\boldsymbol{\xi}}, \quad t \rightarrow \infty, \quad (9)$$

where  $\boldsymbol{\sigma}(\cdot) = (\sigma_j(\cdot))_{j=1}^d$  contains the scaling functions in Assumption 3.

The proof of Proposition 1 is provided in Appendix A.

Typically, as for the univariate case, the scaling function in (9) is viewed as a scaling parameter and Proposition 1 leads to the approximation in distribution for

$$\mathbf{Y}_{\mathbf{u}(t)} \stackrel{\mathcal{L}}{\approx} \mathbf{H} \stackrel{\text{def}}{=} \boldsymbol{\sigma} \frac{\mathbf{Y}^\xi - 1}{\boldsymbol{\xi}}, \quad t \rightarrow \infty \quad (10)$$

with  $\boldsymbol{\sigma} \in (0, \infty)^d$  and  $\boldsymbol{\xi} \in \mathbb{R}^d$  containing, respectively, the scale and shape parameters from the marginal GP distributions in Assumption 3 and where  $\mathbf{Y}$  is Multivariate Pareto distributed as in (8).

*Remark 2* (Standard MGP random vectors). The random vector  $\mathbf{H}$  in (10) is  $\text{MGP}(\boldsymbol{\sigma}, \boldsymbol{\xi}, \mathbf{S})$  distributed, in the sense of [43, Definition 2.2], with *spectral generator*  $\mathbf{S}$  having distribution

$$\mathbb{P}(\mathbf{S} \leq \mathbf{x}) = \frac{\mathbb{E}[\exp(Q) \mathbb{I}\{\mathbf{U} \leq \mathbf{x} + Q\}]}{\mathbb{E}[\exp(Q)]}, \quad \mathbf{x} \in \mathbb{R}^d, \quad (11)$$

where  $\mathbb{I}\{\mathcal{E}\}$  denotes the indicator random variable of the event  $\mathcal{E}$  and  $\mathbf{U} \stackrel{\text{def}}{=} \log(\boldsymbol{\Theta})$ , with  $\boldsymbol{\Theta} \sim \Phi$ , and  $Q \stackrel{\text{def}}{=} \max(\mathbf{U})$ . In other words, the random vector  $\mathbf{U}$  is a  $U$ -generator of  $\mathbf{H}$  in the sense of Proposition 9 of [50]; see also equation (12) in [43]. Note in particular that the moment condition  $0 < \mathbb{E}[\exp(U_j)] < \infty$  for all  $j \in \{1, \dots, d\}$  is satisfied as we have, from (3),

$$\mathbb{E}[\exp(U_j)] = \mathbb{E}[\Theta_j] = \int_{\Delta_{d-1}} w_j \, d\Phi(\mathbf{w}) = \frac{1}{d}, \quad j \in \{1, \dots, d\}.$$

Here, since we work mainly with the angular measure  $\Phi$  in our procedure, we decided to consider the multivariate Pareto random vector  $\mathbf{Y}$ , as in (8), to be the “standard”

MGP, that is with  $\boldsymbol{\xi} = \mathbf{1}$  and  $\boldsymbol{\sigma} = \mathbf{1}$  instead of the choice  $\boldsymbol{\xi} = \mathbf{0}$  made in [50, 43, 51], for example. This approach has also been considered in [23]. Of course, it suffices to take  $\mathbf{Z} = \log(\mathbf{Y})$  to obtain an MGP random vector with parameter  $\boldsymbol{\xi} = \mathbf{0}$ , as can be seen from comparing (10) to Equation (4) in [43].

**Tail modeling and rare event probabilities.** Thanks to Proposition 1, we have now a clear path to modeling the tail of  $\mathbf{X}$ . Suppose we are interested in estimating a probability  $P(\mathbf{X} \in C)$  for some  $C \subseteq \{\mathbf{x} \in \mathbb{R}^d : \mathbf{x} \not\leq \mathbf{u}(t)\}$  for some large  $t > 1$ , with the same  $\mathbf{u}(t)$  as in (9), so that the approximation of  $\mathbf{Y}_{\mathbf{u}(t)} = \mathbf{X} - \mathbf{u}(t) \mid \mathbf{X} \not\leq \mathbf{u}(t)$  by the MGP  $\mathbf{H}$  as in (10) is accurate. In such a region, there may be no or few observations available, so that an empirical estimate is not reliable. Instead, we write the probability of interest as

$$\begin{aligned} P(\mathbf{X} \in C) &= P(\mathbf{X} \not\leq \mathbf{u}(t)) P(\mathbf{Y}_{\mathbf{u}(t)} \in C - \mathbf{u}(t)) \\ &\approx P(\mathbf{X} \not\leq \mathbf{u}(t)) P(\mathbf{H} \in C - \mathbf{u}(t)). \end{aligned} \quad (12)$$

The first probability can be estimated using an empirical evaluation on the sample while the second one will be computed using an estimate of the MGP distribution of  $\mathbf{H}$ .

As can be seen from (10), the MGP random vector  $\mathbf{H}$  is a simple transformation of the standard MGP random vector  $\mathbf{Y}$  introduced in (8), involving only marginal parameters for which estimation is well developed mathematically and numerically. Therefore, the main challenge to obtain informative samples to evaluate tail probabilities related to  $\mathbf{X}$  is to obtain realizations of  $\mathbf{Y}$ . By Proposition 1, this is possible if we are able to simulate from  $\Phi$ —at least approximately since it is a limit distribution by nature, so that no direct observation of it will ever be available.

We propose to use a specific GAN structure to achieve this goal. The main principles underlying this framework are recalled in the following section.

## 2.2 Wasserstein Generative Adversarial Networks

Suppose we are interested in estimating the unknown distribution  $P$  of some, potentially complex, data on a large vector space  $\mathcal{X}$ . For instance, think of the angular measure  $\Phi$  in Section 2.2 supported on the simplex  $\Delta_{d-1}$  for large  $d$ . Then, the standard parametric approaches based on densities may not be satisfactory as they are often too restrictive to take into account all the possible characteristics of a complex angular measure in large dimension. An alternative direction is to start from a simple random variable  $Z \in \mathcal{Z}$  in some “latent” space—typically, of lower dimension than  $\mathcal{X}$ —whose distribution is known and to look for a parametric transformation  $g_\theta : \mathcal{Z} \rightarrow \mathcal{X}$  such that  $g_\theta(Z)$  has a distribution sufficiently close to  $P$ . Practically, this has several advantages. We get more flexibility as the map  $g_\theta$  can be virtually anything if we consider a sufficiently rich family of parametric functions and the latent space/distribution can also be selected freely. It is easier to simulate from  $P$  as it suffices to obtain a sample of  $Z$  (easy) and transform it based on  $g_\theta$  rather than more costly algorithms based on the knowledge of the density values.

A Wasserstein GAN (WGAN) architecture [5] follows this path to estimate  $\mathbf{P}$  by considering  $g_\theta$  to be a neural network and by measuring the difference in distribution between the real data distribution  $\mathbf{P}$  and the law of  $g_\theta(Z)$  via the *Earth Mover's Distance* (EMD), which is a particular instance of the Wasserstein distance in optimal transport [56].

A WGAN consists of two neural networks: a so-called “Generator”  $G = G_\theta : \mathcal{Z} \mapsto \mathcal{X}$  and a “Discriminator”  $D = D_w : \mathcal{X} \mapsto \mathbb{R}$  (also sometimes called “Critic” depending on the authors). These two networks play an opposite role:  $G_\theta$  aims to produce samples as realistic as possible and close to the target distribution and  $D_w$  aims at efficiently discriminating samples and determining which ones are generated by  $G_\theta$  and which ones are real. This “game” can be formulated as a min-max optimization problem in the EMD setup that we recall know.

The EMD is simply a distance between probability measures with finite first moment defined on a metric space  $(\mathcal{X}, \rho)$ . If  $\mathbf{P}$  and  $\mathbf{Q}$  are two such distributions, it is equivalently defined as

$$W_1(\mathbf{P}, \mathbf{Q}) \stackrel{\text{def}}{=} \inf_{\pi \in \Pi(\mathbf{P}, \mathbf{Q})} \int_{\mathcal{X} \times \mathcal{X}} \rho(x, y) \, d\pi(x, y) \quad (13)$$

$$= \sup_{f \in \text{Lip}_1} \left\{ \int_{\mathcal{X}} f(x) \, d\mathbf{P}(x) - \int_{\mathcal{X}} f(x) \, d\mathbf{Q}(x) \right\}, \quad (14)$$

where  $\Pi(\mathbf{P}, \mathbf{Q})$  is the set of probability measures on  $\mathcal{X} \times \mathcal{X}$  with  $\mathbf{P}$  and  $\mathbf{Q}$  as first and second margins and  $\text{Lip}_1$  is the set of functions  $f : \mathcal{X} \rightarrow \mathbb{R}$  that are 1-Lipschitz continuous, that is, functions that satisfy

$$\sup_{x, y \in \mathcal{X}, x \neq y} \frac{|f(x) - f(y)|}{\rho(x, y)} \leq 1.$$

If  $\mathcal{X}$  is an open subset of  $\mathbb{R}^N$  for some  $N \in \mathbb{N}$ , this implies

$$|\nabla f(x)|_2 \leq 1 \quad \text{for almost every } x \in \mathcal{X}, \quad (15)$$

that is, for every  $x \in \mathcal{X}$  except for a subset of Lebesgue measure zero. The equivalence between (13) and (14) is known as the Kantorovich–Rubinstein duality theorem, the proof of which is to be found in, e.g., [56, Section 1.2].

In formulation (14) of the EMD, the discriminator  $D_w$  will play the role of the “separating” function  $f$  and  $\mathbf{Q}$  will correspond to the measure generated by the generator  $G_\theta$ , that is, the distribution of  $G_\theta(Z)$ , while  $\mathbf{P}$  will again be our target measure from which we observe a random sample. This leads to the following min-max procedure:

$$\min_{\theta \in \mathcal{S}_G} \max_{w \in \mathcal{S}_D} \left\{ \int_{\mathcal{X}} D_w(x) \, d\mathbf{P}(x) - \int_{\mathcal{Z}} D_w(G_\theta(z)) \, d\mathbf{P}_Z(z) \right\}, \quad (16)$$

where  $\mathbf{P}_Z$  denotes the law of the latent vector  $Z$  and  $\mathcal{S}_G$  and  $\mathcal{S}_D$  denote the (finite-dimensional) parameter spaces of the generator and the discriminator respectively. Optimization is performed in the respective parameter spaces of the generator and the



discriminator, which contain all the weights in their respective network structures. Typically, the expectations are replaced by averages over batches of real data with distribution  $\mathbf{P}$  and fake data with distribution  $\mathbf{P}_Z$ . The joint minimization/maximization of the objective function is performed by gradient descent based on the standard backpropagation algorithm. The whole GAN procedure is summarized in Figure 1.

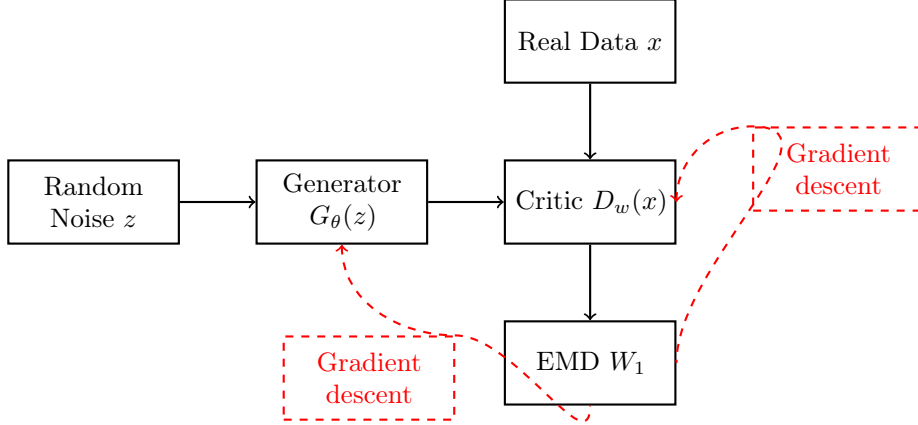


Figure 1: Illustration of a Wasserstein GAN with gradient-based optimization

In practice there is no need for the neural network  $D_w$  to be Lipschitz continuous as required in definition (14) of the EMD. Initially, in [5], the authors proposed to clip the weights, that is, forcing them to lie in a fixed compact space, typically  $[-c, c]^{|\mathcal{S}_D|}$  for some small  $c > 0$ . One may show that if the discriminator consists of a multilayer perceptron, which will always be the case here, its Lipschitz norm will be bounded by a finite constant depending only on  $c$ . Even though this constant may not be equal to one, replacing  $\text{Lip}_1$  by  $\text{Lip}_C$  for  $C > 0$  in (14) only amounts to multiplying the value of the resulting distance by  $C > 0$  and it does not play a role in the optimization of the generator.

Since then, it has been shown that in some cases it is preferable to enforce the Lipschitz constraint in a soft manner, using (15). Of course, controlling the gradient  $\nabla D_w$  everywhere on  $\mathcal{X}$  is not possible, so what is usually done is to add a penalty term to the objective function of the form

$$\lambda \int_{\mathcal{X}} (|\nabla D_w(x)|_2 - 1)^2 d\mu(x), \quad \lambda > 0, \quad (17)$$

where  $\mu$  is a probability measure on  $\mathcal{X}$ . The standard approach in [27] is to take  $\mu$  as the law of the mixture  $U \cdot X_r + (1 - U) \cdot G_\theta(Z)$  where  $X_r \sim \mathbf{P}$  follows the data distribution,  $Z \sim \mathbf{P}_Z$  is random noise and  $U$  is uniformly distributed on  $[0, 1]$  and independent of the rest. This choice is motivated by Proposition 1 in [27] and we follow this approach here.

As explained at the end of Section 2.1 and in the beginning of this section, our aim is to use this architecture to obtain accurate samples from the angular measure  $\Phi$  on  $\Delta_{d-1}$ . Since the simplex structure of such observations makes their distribution a bit non-standard

and could prevent the generator from producing high-quality output, we propose to apply the procedure on a transformation of the angular data to coordinates in an orthonormal basis of the simplex and to transform the coordinates back to angles afterwards. The next section and Appendix B provide more details about this transformation.

### 2.3 Aitchison coordinates

A core idea of compositional data analysis [1, 18], mainly developed by John Aitchison (1926–2016), consists of transforming raw simplex data in  $\mathbb{R}^d$  to coordinates with respect to an orthonormal basis of a certain  $(d - 1)$ -dimensional subspace of  $\mathbb{R}^d$ . This makes interpretation more complicated, but having data on a full linear  $(d - 1)$ -dimensional sample space instead of on the unit simplex in  $\mathbb{R}^d$  can improve statistical and neural network modeling substantially. Since we are not all that concerned with interpretation in the WGAN architecture described in Section 2.2 but instead with obtaining the highest possible efficiency in producing accurate new multivariate extreme samples, we will apply this idea in our context.

Infinitely many orthogonal bases exist for the open simplex  $\Delta_{d-1}^o$  equipped with its inner product space structure. Here, we work with the one proposed in Proposition 2 of [18]. Details of the underlying construction and more information on the Aitchison simplex, i.e., the unit simplex equipped with a certain vector space structure and an inner product, can be found in Appendix B. The essence of the matter is to map  $\Delta_{d-1}^o$  to  $\mathbb{H} = \{\mathbf{x} \in \mathbb{R}^d : \sum_{j=1}^d x_j = 0\}$  by means of the centered logratio function (clr) in (28). The set  $\mathbb{H}$  is a  $(d - 1)$ -dimensional linear space equipped with the standard Euclidean inner product, the structure of which carries over to  $\Delta_{d-1}^o$  through the function clr. Recall that for a vector  $\mathbf{x} \in \mathbb{R}^d$ ,  $\text{softmax}(\mathbf{x})$  is the element of  $\Delta_{d-1}^o$  defined by

$$\text{softmax}(\mathbf{x}) \stackrel{\text{def}}{=} \left( \frac{\exp(x_j)}{\sum_{i=1}^d \exp(x_i)} \right)_{j=1}^d.$$

**Proposition 2** (An orthonormal basis for the Aitchison simplex). *The vectors  $\mathbf{e}_1^*, \dots, \mathbf{e}_{d-1}^*$  defined by*

$$\mathbf{e}_i^* \stackrel{\text{def}}{=} \text{softmax}(e_i) \in \Delta_{d-1}^o \text{ with } e_i \stackrel{\text{def}}{=} \sqrt{\frac{i}{i+1}} \left( \underbrace{i^{-1}, \dots, i^{-1}}_{i \text{ times}}, -1, 0, \dots, 0 \right) \in \mathbb{R}^d$$

for  $i \in \{1, \dots, d - 1\}$  form an orthonormal basis of the Aitchison simplex.

Now that we have constructed an orthonormal basis of the Aitchison simplex in Proposition 2, we are able to decompose any vector  $\mathbf{u} \in \Delta_{d-1}^o$  in this basis as

$$\mathbf{u} = \bigoplus_{i=1}^{d-1} \langle \mathbf{u}, \mathbf{e}_i^* \rangle_A \odot \mathbf{e}_i^* = \bigoplus_{i=1}^{d-1} \langle \text{clr}(\mathbf{u}), \text{clr}(\mathbf{e}_i^*) \rangle \odot \mathbf{e}_i^* = \bigoplus_{i=1}^{d-1} \langle \text{clr}(\mathbf{u}), \mathbf{e}_i \rangle \odot \mathbf{e}_i^*.$$

the notation being defined in Appendix B. The coordinates  $\mathbf{u}_i^* \stackrel{\text{def}}{=} \langle \text{clr}(\mathbf{u}), \mathbf{e}_i \rangle$  for  $i \in \{1, \dots, d-1\}$  are sufficient to completely describe the vector  $\mathbf{u}$ . These coordinates will be the quantities entering the WGAN structure described in Section 2.2 as they typically have a nicer distribution than the original simplex random vectors of interest and hence will be easier to “learn”, which helps facing the difficulty that sample sizes in extreme value analysis are often not very large.

### 3 Combining GAN, Aitchison coordinates, and Extremes

Recall that our main objective is to provide an algorithm which is able to accurately sample from the tail of a random vector  $\mathbf{X} \in \mathbb{R}^d$ . We do this by combining Algorithm 1 and Algorithm 2 described below.

- Algorithm 1 develops a Wasserstein GAN to simulate from the dependence structure of the unit-Pareto transformed observations.
- Algorithm 2 uses estimates of the marginal GP parameters to transform these simulated values to the original scale.

The simulated values can then be used to obtain accurate estimates of probabilities of the form  $\mathbf{P}(\mathbf{X} \in C)$  where  $C \subseteq \{\mathbf{x} \in \mathbb{E} : \mathbf{x} \not\leq \mathbf{u}(t)\}$  for some “high” threshold vector  $\mathbf{u}(t) \in \mathbb{R}^d$  as in (9), close to  $\mathbf{u}_*$ .

The main challenge lies in simulation from the angular measure  $\Phi$  of  $\mathbf{V}$  in Assumption 1, since moving from  $\Phi$  to the MGP distribution is a matter of scaling and marginal tail estimation, for which theory and practice are well developed. Consequently, we propose to use the WGAN architecture of Section 2.2 to obtain accurate samples from  $\Phi$ . This is done via the Aitchison coordinates, Section 2.3, as detailed in Algorithm 1 below.

The input to Algorithm 1 is a set of training observations  $\mathbf{X}_1, \dots, \mathbf{X}_n$ , with  $\mathbf{X}_i = (X_{i1}, \dots, X_{id})$ , which will be converted to approximately unit-Pareto margins using the empirical marginal cumulative distribution functions

$$\hat{F}_j(x) \stackrel{\text{def}}{=} \frac{1}{n+1} \sum_{i=1}^n \mathbb{I}\{X_{ij} \leq x\}, \quad x \in \mathbb{R}, j \in \{1, \dots, d\}. \quad (18)$$

Division by  $n+1$  instead of  $n$  is there to prevent division by zero in the standardization. The angles associated to “large” observations are supposed to provide an approximate sample from  $\Phi$  in view of Assumption 1. The output of Algorithm 1 is a trained generator  $G_\theta$  which is able to produce new coordinates in the Aitchison orthonormal basis, defined by the orthonormal basis  $\{\mathbf{e}_1, \dots, \mathbf{e}_{d-1}\}$  of  $\mathbb{H}$ , that are hopefully close in distribution to the coordinates of angles associated to the large observations in the training set.

The WA-GAN procedure in Algorithm 1 only uses the observations for which  $R = \|\mathbf{V}\|_1$  is larger than  $t = n/k_1$  so that, by (4), the number of observations used by the algorithm is asymptotically binomial with parameters  $n$  and  $(dk_1)/n$ . Here  $k_1$  is a value to be chosen by the user, with larger  $k_1$  meaning that more observations are used and hence

the variance is smaller, and with smaller  $k_1$  (hopefully) making the approximation in Assumption 1 more accurate. Finite sample bounds on the quality of this approximation are given in [12]. The relatively low effective sample size makes it important that the architectures of the generator and the discriminator are simple enough so that their parameters can be learned from a moderate number of examples. Those considerations are further explored in the numerical sections.

Algorithm 2 uses for each margin the  $k_2$  largest observations, so that between  $k_2$  and  $dk_2$  of the multivariate observations are used. According to Remark 4, the value of  $k_2$  should not be larger than  $k_1$ . A natural convenient choice is to take  $k_1 = k_2$ .

Algorithm 1 is essentially the standard WGAN algorithm with gradient penalty of [27] applied to the coordinates of the large angles in an orthonormal basis of the Aitchison simplex. The ADAM optimizer which is used to update the weights of the generator and the discriminator at each step is a popular gradient descent algorithm with adaptive learning rate introduced in [35]. The function `Adam()` in the algorithm is one such gradient step and is the default optimizer proposed for the WGAN architecture with gradient penalty. Compared to the standard algorithm, one additional regularization term

$$\rho \left\| \frac{1}{m} \sum_{i=1}^m \text{softmax}(MG_\theta(\mathbf{z}_i)) - \frac{\mathbf{1}_d}{d} \right\|_2^2, \quad (19)$$

where  $M$  is the matrix defined in Algorithm 1, is added to enforce the marginal constraint (3) in the generator in a soft way and thus to enforce sampling from a genuine angular measure. The expression in (19) penalizes the (empirical) deviations from the marginal conditions in (3) by forcing the squared Euclidean norm of their differences to be close to zero. Note that  $MG_\theta(\mathbf{z}_i)$  is the point in  $\mathbb{H}$  with coordinates  $G_\theta(\mathbf{z}_i)$  with respect to the orthonormal basis  $\{\mathbf{e}_1, \dots, \mathbf{e}_{d-1}\}$ ; the softmax function maps this point to the unit simplex. The regularization term then encourages the empirical distribution of the  $m$  points on the unit simplex obtained in this way to the moment constraints in (3).

Algorithm 1 focuses solely on the extremal dependence of  $\mathbf{X}$  and is independent of any marginal assumptions besides continuity of the marginal distribution functions. Therefore, if one is interested in computing a quantity which only depends on  $\Phi$  in Assumption 1, such as the coefficients  $\theta_j$  introduced in the beginning of Section 4, this can be done independently of Algorithm 2. This is not the case for methods such as [37] which directly evaluate sample quality on the scale of the MGP distribution. Note also that even though Algorithm 1 only considers the angular measure  $\Phi$  with respect to the  $L_1$ -norm, a simple post-processing step allows for any norm on  $\mathbb{R}^d$ , as explained in Appendix C.

*Remark 3* (Standardization of the margins to unit-Pareto). The procedure in Algorithm 1 is based on data standardized to unit-Pareto margins, in accordance with Assumption 1. As in (18), we work with an empirical standardization  $\hat{F}_j$  for each margin  $j = 1, \dots, d$ , resulting in a procedure using only the marginal ranks of the original data and solely focusing on dependence, not requiring any assumption on the tail of  $F_j$  as in Assumption 3. The same approach was followed in [19, 21] and, more recently, in [12], for example. An alternative would be to use the empirical distribution function  $\hat{F}_j$  only up to a threshold

---

**Algorithm 1** WA-GAN for tail dependence estimation
 

---

**Require:** observations  $\mathbf{X}_1, \dots, \mathbf{X}_n$ , orthonormal basis  $\{\mathbf{e}_1, \dots, \mathbf{e}_{d-1}\}$  of  $\mathbb{H}$ , and  $k_1 \in \{1, \dots, n\}$

**Require:** gradient penalty coefficient  $\lambda > 0$ , marginal penalty coefficient  $\rho > 0$ , the number of discriminator iterations per generator iteration  $n_D$ , the batch size  $m$ , Adam hyperparameters  $\alpha > 0$  (learning rate),  $\beta_1, \beta_2 \in [0, 1)$ , the latent space dimension  $\ell \in \mathbb{N}$ , the number of epochs  $n_E$

**Require:** discriminator initial parameters  $w_0 \in \mathcal{S}_D$ , generator initial parameters  $\theta_0 \in \mathcal{S}_G$

$t \leftarrow n/k_1$

**for**  $i = 1, \dots, n$  **do**

$\hat{\mathbf{V}}_i \leftarrow \left( 1/(1 - \hat{F}_j(X_{ij})) \right)_{j=1}^d$

$R_i \leftarrow \|\hat{\mathbf{V}}_i\|_1$  and  $\mathbf{W}_i \leftarrow \hat{\mathbf{V}}_i/R_i$

**if**  $R_i \geq t$  **then**

$\mathbf{W}_i^* \leftarrow (\langle \text{clr}(\mathbf{W}_i), \mathbf{e}_j \rangle)_{j=1}^{d-1}$

**end if**

**end for**

$K \leftarrow \sum_{i=1}^n \mathbb{I}\{R_i \geq t\}$

$M \leftarrow (\mathbf{e}_1 \mid \dots \mid \mathbf{e}_{d-1}) \in \mathbb{R}^{d \times (d-1)}$

$\theta \leftarrow \theta_0$  and  $w \leftarrow w_0$

**for** epoch  $e = 1, \dots, n_E$  **do**

**for**  $t = 1, \dots, n_D$  **do**

**for**  $i = 1, \dots, m$  **do**

sample  $\mathbf{w}_i^*$  from  $\{\mathbf{W}_1^*, \dots, \mathbf{W}_K^*\}$ , sample  $\mathbf{z}_i$  from a standard multivariate normal distribution on  $\mathbb{R}^\ell$ , sample  $u_i$  from a uniform distribution on  $[0, 1]$

$\tilde{\mathbf{w}}_i^* \leftarrow G_\theta(\mathbf{z}_i)$

$\hat{\mathbf{w}}_i^* \leftarrow u_i \mathbf{w}_i^* + (1 - u_i) \tilde{\mathbf{w}}_i^*$

**end for**

$L_D \leftarrow \frac{1}{m} \sum_{i=1}^m \left\{ D_w(\tilde{\mathbf{w}}_i^*) - D_w(\mathbf{w}_i^*) + \lambda (|\nabla_{\mathbf{w}^*} D_w(\hat{\mathbf{w}}_i^*)|_2 - 1)^2 \right\}$

$w \leftarrow \text{Adam}(L_D, w, \alpha, \beta_1, \beta_2)$

**end for**

**for**  $i = 1, \dots, m$  **do**

sample  $\mathbf{z}_i$  from a standard multivariate normal distribution on  $\mathbb{R}^\ell$

**end for**

$L_G \leftarrow -\frac{1}{m} \sum_{i=1}^m D_w(G_\theta(\mathbf{z}_i)) + \rho \left\| \frac{1}{m} \sum_{i=1}^m \text{softmax}(MG_\theta(\mathbf{z}_i)) - \frac{\mathbf{1}_d}{d} \right\|_2^2$

$\theta \leftarrow \text{Adam}(L_G, \theta, \alpha, \beta_1, \beta_2)$

**end for**

**Output:** Trained generator  $G_\theta$

---

$u_j < u_{*j}$  and fit a univariate GP distribution above this threshold, relying on Assumption 3. Doing so could improve the quality of the transformation to unit-Pareto margins, but at the price of making Algorithm 1 also require marginal assumptions. Furthermore, this route has already been explored in [20] and the same authors advocate in [19] for the pure nonparametric approach. Therefore, we stick to the rank-based procedure.

Some care is needed when sampling from  $\mathbf{X} \mid \mathbf{X} \not\leq \mathbf{u}(t)$  for  $t > 1$  large, where  $\mathbf{u}(t)$  is the vector of thresholds such that for each component separately, the excess probability is  $1/t$ . For example, a common situation is the case where  $\mathbf{X}$  represents some positive multivariate risks, that is,  $\mathbf{X}$  takes values in  $[0, \infty)^d$ . If we directly make use of the formula  $\mathbf{u}(t) + \mathbf{H}$ , with estimated quantities, as suggested by formula (10), we risk generating “extreme” points for which some coordinates are negative. Our approach avoids this issue.

---

**Algorithm 2** Sampling from the tail of  $\mathbf{X}$

---

**Require:** observations  $\mathbf{X}_1, \dots, \mathbf{X}_n$ , orthonormal basis  $\{\mathbf{e}_1, \dots, \mathbf{e}_{d-1}\}$  of  $\mathbb{H}$ , trained generator  $G_\theta$ , the size of the latent space on which  $G_\theta$  has been trained  $\ell$ , and  $k_2 \in \{1, \dots, n\}$

**Require:** desired number of samples  $n^*$

**for**  $j = 1, \dots, d$  **do**

$u_j \leftarrow X_{n-k_2:n,j}$

Compute the maximum likelihood estimator  $(\hat{\sigma}_j, \hat{\xi}_j)$  of a univariate GP distribution

**end for**

$M \leftarrow (\mathbf{e}_1 \mid \dots \mid \mathbf{e}_{d-1}) \in \mathbb{R}^{d \times (d-1)}$

$i \leftarrow 1$  and  $\mathcal{G}$  an empty list of size  $n^*$

**while**  $i \leq n^*$  **do**

sample  $\mathbf{z}$  from a standard multivariate normal distribution on  $\mathbb{R}^\ell$

$\mathbf{W} \leftarrow \text{clr}^{-1}(MG_\theta(\mathbf{z}))$

sample  $Y$  from a unit-Pareto distribution

$\mathbf{Y} \leftarrow Y\mathbf{W}$

**if**  $\max(\mathbf{Y}) > 1$  **then**

**for**  $j = 1, \dots, d$  **do**

**if**  $Y_j \leq 1$  **then**

$(\mathcal{G}[i])_j \leftarrow X_{([n-k_2/Y_j] \vee 1):n,j}$

**end if**

**if**  $Y_j > 1$  **then**

$(\mathcal{G}[i])_j \leftarrow u_j + \hat{\sigma}_j(Y_j^{\hat{\xi}_j} - 1)/\hat{\xi}_j$

**end if**

**end for**

$i \leftarrow i + 1$

**end if**

**end while**

**Output:** List of  $n^*$  random variables  $\mathcal{G}$  approximately distributed as  $\mathbf{X} \mid \mathbf{X} \not\leq \mathbf{u}$

---

Algorithm 1 permits to sample from the multivariate Pareto distribution  $\mathbf{Y}$  associated

to  $\mathbf{X}$  using (8). Let  $\mathbf{Y}^*$  denote such a generated quantity. From (25), we have

$$\mathbf{Y}^* \stackrel{\mathcal{L}}{\approx} \left( \frac{k}{n} \mathbf{V} \mid \mathbf{V} \not\leq \frac{n}{k} \right),$$

where  $\stackrel{\mathcal{L}}{\approx}$  means an approximation in distribution. Equivalently,

$$\mathbf{V}^* \stackrel{\text{def}}{=} \frac{n}{k} \mathbf{Y}^* \stackrel{\mathcal{L}}{\approx} \left( \mathbf{V} \mid \mathbf{V} \not\leq \frac{n}{k} \right).$$

Now, to pass from  $\mathbf{V}$  to  $\mathbf{X}$  one needs to apply the transformation  $\mathbf{b}$  described in Proposition 1, which is based on the marginal quantile functions. Let  $\hat{\mathbf{b}}$  denote an estimator of this quantity. The final output of our next algorithm will be

$$\mathbf{X}^* = \hat{\mathbf{b}}(\mathbf{V}^*) = \hat{\mathbf{b}}\left(\frac{n}{k} \mathbf{Y}^*\right) \stackrel{\mathcal{L}}{\approx} \left( \mathbf{X} \mid \mathbf{V} \not\leq \frac{n}{k} \right) = \left( \mathbf{X} \mid \mathbf{X} \not\leq \mathbf{b}\left(\frac{n}{k}\right) \right),$$

by the continuity of margins. The estimator  $\hat{b}_j((n/k)y)$  for  $y > 0$  depends on whether  $y \leq 1$  or  $y > 1$ . For  $y \leq 1$  we simply use the estimator based on the empirical cumulative distribution function  $\hat{F}_j$ . For  $y > 1$ , the estimator is obtained through Assumption 3, more precisely its equivalent formulation (ii) in Remark 1. This leads to

$$\hat{b}_j((n/k)y) \stackrel{\text{def}}{=} \begin{cases} X_{(\lceil n-k/y \rceil \vee 1):n,j}, & 0 < y \leq 1, \\ \hat{b}_j(n/k) + \hat{a}_j(n/k) \cdot \frac{y^{\hat{\xi}_j} - 1}{\hat{\xi}_j}, & y > 1, \end{cases} \quad (20)$$

where  $X_{1:n,j} < \dots < X_{n:n,j}$  are the ascending order statistics in the  $j$ -th sample, and where  $\hat{a}_j(n/k) \stackrel{\text{def}}{=} \hat{\sigma}_j$  and  $\hat{\xi}_j$  are obtained by maximum likelihood for the GP distribution. *Remark 4* (Relation between thresholds in both algorithms). When combining Algorithm 2 with the output of Algorithm 1, one actually uses two different thresholds: on the one hand, the threshold  $t_1 = n/k_1 > 0$  which has been used to fit the angular measure in Algorithm 1, and, on the other hand, the threshold vector  $\mathbf{u}(t_2) = \mathbf{b}(n/k_2)$  for which we aim to sample from  $\mathbf{X} \mid \mathbf{X} \not\leq \mathbf{u}(t_2)$ . As Algorithm 2 relies on the output of Algorithm 1, one should ensure that  $\mathbf{u}(t_2)$  is large enough so that the approximation of the angular measure  $\Phi$  in Algorithm 1 is valid. In terms of the random variable  $R$  in (1), it means that we must ensure  $R = |\mathbf{V}|_1 > n/k_1$  given that  $\mathbf{X} \not\leq \mathbf{u}(t_2)$ , which is equivalent to  $\mathbf{V} \not\leq n/k_2$ , i.e.,  $|\mathbf{V}|_\infty > n/k_2$ . Since we have  $|\mathbf{V}|_1 \geq |\mathbf{V}|_\infty$ , it suffices to take  $k_2 \leq k_1$ . This is illustrated in Figure 2 for  $d = 2$ .

## 4 Numerical experiments

We use both simulated and real data sets to explore the performance of our proposed method and to compare it to other existing methods. We start with providing the details of how the different methods are compared and how the computations are made. The next section uses simulated data to compare WA-GAN with HTGAN and GPGAN. In the final section these methods are applied to a financial data set.

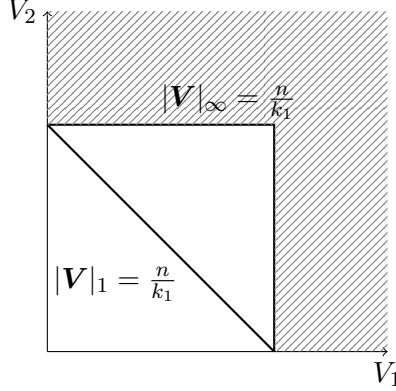


Figure 2: Relation between thresholds in the  $\mathbf{V}$  space. The gray zone is where we can safely sample using Algorithm 2.

#### 4.1 Performance and computation

**Performance measures.** To assess the quality of the generated sample  $\mathbf{X}^{\mathcal{G}}$ , we will typically compare it with a test set  $\mathbf{X}^{\mathcal{T}}$ , which is a part of the data that has not been used for training. We propose to use measures that are inspired by multivariate extreme value theory since our focus is on the tail of  $\mathbf{X}$ .

A first set of measures that we propose is based on the *extremal coefficients* [7, Section 8.2.7]. For any non-empty subset  $J \subseteq \{1, \dots, d\}$  with at least two elements  $|J| \geq 2$ , we define its extremal coefficient  $\theta_J$  as

$$\theta_J \stackrel{\text{def}}{=} d \cdot \int_{\Delta_{d-1}} \bigvee_{j \in J} w_j \, d\Phi(\mathbf{w}) \in [1, |J|]. \quad (21)$$

The closer  $\theta_J$  is to 1, the stronger the dependence in the tail of  $(X_j)_{j \in J}$ , while  $\theta_J$  close to  $|J|$  corresponds to weaker tail dependence. For a given coefficient order  $k = |J|$ , to compute a summary score of their difference in the test set  $\mathbf{X}^{\mathcal{T}}$  and the generated set  $\mathbf{X}^{\mathcal{G}}$ , we compute their relative absolute difference as follows:

$$E_{\bar{\theta}}(k) \stackrel{\text{def}}{=} \frac{1}{\binom{d}{k}} \sum_{\substack{J \subseteq \{1, \dots, d\} \\ |J|=k}} \left| 1 - \frac{\theta_J^{\mathcal{G}}}{\theta_J^{\mathcal{T}}} \right|. \quad (22)$$

The coefficients are estimated for the test set using the empirical angular measure in (21) and for the generative methods using the sampled realizations of  $\Phi$ . For the test set, the radial threshold  $n/k$  is the same as for the training set. Typically, we will consider the bivariate coefficients, with  $k = 2$ , as in [11], and the trivariate coefficients, with  $k = 3$ . This measure focuses on the extremal dependence structure only and aims to measure the quality of Algorithm 1.

A second measure that we propose, which also focuses on the tail of  $\mathbf{X}$ , aims at evaluating the output of Algorithm 2. Let  $\mathbf{X}^{\mathcal{T}}$  denote a set of observations distributed



like  $\mathbf{X}$ , which typically have not been used for training, and let  $\mathbf{X}_u^{\mathcal{G}}$  denote a set of observations generated by Algorithm 2 based on the threshold vector  $\mathbf{u} = \mathbf{b}(n/k_n) \in \mathbb{R}^d$ . Then if the MGP approximation is accurate enough, we expect  $\mathbf{X}_u^{\mathcal{T}} \stackrel{\text{def}}{=} \mathbf{X}^{\mathcal{T}} | \mathbf{X}^{\mathcal{T}} \not\leq \mathbf{u}$  to be close in distribution to  $\mathbf{X}_u^{\mathcal{G}}$ . Hence, we propose to compute their difference in distribution using a simple extension of the EMD in Section 2.2 which is known as the 2-Wasserstein distance [56, Chapter 7]. If  $n_{\mathcal{G}}$  and  $n_{\mathcal{T}}$  denote, respectively, the size of the test set and the generated set of observations, the measure can be expressed as

$$W_2(\mathbf{X}_u^{\mathcal{G}}, \mathbf{X}_u^{\mathcal{T}}) \stackrel{\text{def}}{=} \inf_{\pi \in \Pi\left(\frac{\mathbf{1}_{n_{\mathcal{G}}}}{n_{\mathcal{G}}}, \frac{\mathbf{1}_{n_{\mathcal{T}}}}{n_{\mathcal{T}}}\right)} \sqrt{\sum_{i=1}^{n_{\mathcal{G}}} \sum_{j=1}^{n_{\mathcal{T}}} \|(\mathbf{X}_u^{\mathcal{G}})_i - (\mathbf{X}_u^{\mathcal{T}})_j\|_2^2 \pi_{ij}} \quad (23)$$

where, for (probability) vectors  $\mathbf{a} \in \mathbb{R}^k$  and  $\mathbf{b} \in \mathbb{R}^{\ell}$ , we use the notation

$$\Pi(\mathbf{a}, \mathbf{b}) \stackrel{\text{def}}{=} \left\{ \pi \in \mathbb{R}_+^{k \times \ell} : \pi \mathbf{1}_{\ell} = \mathbf{a} \text{ and } \pi^T \mathbf{1}_k = \mathbf{b} \right\}.$$

We can view  $\pi_{ij}$  as the amount of mass transferred from  $(\mathbf{X}_u^{\mathcal{G}})_i$  to  $(\mathbf{X}_u^{\mathcal{T}})_j$ . In a similar way as in definition (13) of the EMD,  $W_2(\mathbf{X}_u^{\mathcal{G}}, \mathbf{X}_u^{\mathcal{T}})$  measures, among all the possible ways to transport the empirical distribution of the vectors in  $\mathbf{X}_u^{\mathcal{G}}$  to the one of the vectors in  $\mathbf{X}_u^{\mathcal{T}}$ , the cheapest way to do so if the cost of unitary transportation is equal to the squared Euclidean distance between points. In the case where we would have normally distributed quantities, this quantity would reduce to computing the well-known *Fréchet inception distance*, which is popularly used to assess the quality of generated images [30] using generative systems like the GAN, also used in [37] to assess the quality of the MGP output of their algorithm. As we compare quantities that are far from being normally distributed, we think it makes more sense to compute the 2-Wasserstein distance instead.

**Architectures.** For both the generator and the discriminator in Algorithm 1, we consider fully connected multilayer perceptrons with leaky rectified linear unit activation function [38] which applies the map

$$x \in \mathbb{R} \mapsto \begin{cases} x, & \text{if } x \geq 0, \\ \alpha x, & \text{otherwise,} \end{cases} \quad (24)$$

to each component of the linear transformation in a given layer, where  $\alpha \in [0, 1]$  is typically set to 0.01. Taking  $\alpha = 0$  leads to the rectified linear unit activation function and  $\alpha = 1$  leads to a simple linear layer. We refer to [38] for a comparison with other standard activation functions. Here, we consider  $\alpha = 0.01$  for any hidden layer while we take  $\alpha = 1$  for the final layer in each network. More recently,  $\alpha$  has been considered as a learnable parameter in the optimization process [29] but this option is not considered in this work. Other architectures are possible if the data have some structure. For example, the authors of [11] consider convolutions in their network as they work with spatial data on a grid, and this could be adapted to our framework too, but here we only consider basic multivariate data and keep this extension for future work.

**Hyperparameter search.** In addition to the learnable parameters of the generator and discriminator networks in Algorithm 1, many external parameters (hyperparameters) can be chosen by the user. We consider the following values for the hyperparameters:

- batch size  $m \in \{\lfloor n/k \rfloor : k = 1, 2, 4, 8, 16\}$ ;
- dimension of the generator’s latent space  $\ell \in \{\lfloor (d-1) \cdot k \rfloor : k = 0.25, 0.5, 0.75, 1, 2\}$ ;
- maximum number of neurons in a hidden layer in  $\{32, 64, 128, 256, 512\}$ ;
- number of hidden layers in  $\{1, 2, 4, 8\}$ ;
- learning rate for the ADAM optimizer in  $\{0.01, 0.005, 0.001, 0.0005, 0.0001\}$  and coefficients  $\beta = (\beta_1, \beta_2)$  used for computing running averages of gradient and its square in  $\{(0.0, 0.9), (0.5, 0.9), (0.5, 0.99)\}$ ;
- the gradient penalty coefficient  $\lambda \in \{0.1, 1, 3, 5, 7, 9\}$ ;
- the marginal penalty coefficient  $\rho \in \{0.001, 0.01, 0.1, 1, 3\}$ ;
- the number of discriminator iterations per generator iterations  $n_D \in \{1, 3, 5, 10\}$ .

Inside the space of possible values for those hyperparameters, we perform a random search. Typically, we consider around 2000 different models in this space for a given dataset. We train them with  $n_e = 5000$  epochs, since training them even longer did not significantly improve the performance, but greatly affects the computation time. Evaluation of the associated models is based on the extremal coefficient metric  $E_{\bar{\theta}}(k)$  in (22) with  $k = 2$  and  $k = 3$ . The validated model is chosen to be the one with the lowest value of  $(E_{\bar{\theta}}(2) + E_{\bar{\theta}}(3)) / 2$  among the 2000 models. This validation is performed on a separate dataset from the one used for training or testing.

**Comparison with existing methods.** As reviewed in Section 1, numerous methods already exist in the rapidly growing field that combines generative approaches with extreme value theory, each with its own advantages and limitations. Therefore, an exhaustive comparison with all existing methods is not feasible, even less so as the software underlying the numerical results is often not available from the papers. Instead, we choose to compare our approach with versions of two recent methods: HTGAN from [25] and GPGAN from [37]. We discuss our specific implementations in Appendix D in the supplement.

**Software.** The code used to produce the results of this section was written in **Python**—in particular, the **PyTorch** [44] framework was used to train the generative models—and in **R** [46], which was mainly used for data preprocessing and computation of the performance measures associated with the generated data, notably via the **transport** [53] package for the computation of the Wasserstein distance  $W_2$  in (23).

**Hardware.** The training of the various algorithms and the computation of performance metrics have been performed on the CPUs of the Lemaitre4 and the NIC5 computer clusters, which are managed by the Consortium of Équipements de Calcul Intensif (CÉCI). Details and more specifications related to these computers can be found on the CECI website <https://www.cec-hpc.be/>.

## 4.2 Simulated data

As an illustration of the method’s performance in a controlled setting, we begin by considering simulated data. In Section 4.2.1, we adopt the same experimental setup as in Section 3.2 of [25], with which we compare our method. The dependence structure is logistic, characterized by a single parameter that controls the strength of dependence. This scenario is particularly simple in the sense that each sub-vector exhibits the same form of dependence, and increasing the dimension does not necessarily make the problem more difficult, since the algorithm can easily learn from this highly symmetric structure. For this reason, we consider in Section 4.2.2 a Hüsler–Reiss dependence structure, where each of the  $d(d-1)/2$  pairs of random variables in  $\mathbf{X}$  may have a different level of dependence. Finally, in Section 4.2.3, we turn to an asymptotically independent copula—the Gaussian copula—as the dependence structure of  $\mathbf{X}$ . In this case, Assumption 2 is no longer satisfied, since the angular measure  $\Phi$  is concentrated on the vertices of the simplex. Here, our interest lies in assessing whether our methodology can be improved by allowing for a heavier-tailed latent distribution in the generator.

### 4.2.1 Logistic dependence structure

In this setup, the random vector  $\mathbf{X}$  is generated using a  $d$ -dimensional Gumbel copula with parameter  $\theta \in [1, \infty)$  and Pareto distributed marginals with parameter  $\alpha > 0$ . It is easy to verify that this data-generating process satisfies Assumptions 1 to 3 in Section 2.1. In particular, the extremal coefficient introduced as a performance measure at the beginning of Section 4 is  $\theta_J = |J|^{1/\theta}$ . Hence,  $\theta \rightarrow 1$  corresponds to less dependence in the tail  $\mathbf{X}$ , while  $\theta \rightarrow \infty$  corresponds to perfect tail dependence.

We consider various scenarios in which we assess the performance of WA-GAN and compare it to HTGAN and GPGAN. For HTGAN we follow [25] and make the unrealistic assumption that  $\alpha$  is known. We take  $d \in \{10, 20, 50\}$  and  $\theta \in \{4/3, 2, 4\}$  leading to a Kendall’s tau of  $\tau \in \{1/4, 1/2, 3/4\}$ . For the margins, we consider  $\alpha = 2$  in any scenario. The models are trained on  $n_{\text{train}} = 10\,000$  observations and validated on  $n_{\text{val}} = 5\,000$  observations. The performance assessment is done on  $n_{\text{test}} = 20\,000$  points.

**Results.** For each method and each scenario, two scores are computed. On the one hand, a dependence score equal to  $\{E_{\bar{\theta}}(2) + E_{\bar{\theta}}(3)\}/2$ , with  $E_{\bar{\theta}}(k)$  as in (22), is used to validate the hyperparameters and to assess the performance of the tail dependence in the generated data. On the other hand, to assess the quality of generated extremes  $\mathbf{X} \mid \mathbf{X} \not\leq \mathbf{u}$ , the 2-Wasserstein distance as defined in (23) is computed, where  $u_j = \hat{b}_j(n/k_2)$  with  $\hat{b}_j$  as

in (20), for some  $k_2 \in \{1, \dots, n\}$  which satisfies  $k_2 \leq k_1$  where  $k_1 \in \{1, \dots, n\}$  is used to train the generator in Algorithm 1. Here we take  $n = n_{\text{train}}$  and  $k_1 = k_2 = \lfloor \sqrt{n} \rfloor$ .

The results are reported in Table 1 below.

| $\tau$               | Model  | Dependence score |               |               | Extremes score |               |               |
|----------------------|--------|------------------|---------------|---------------|----------------|---------------|---------------|
|                      |        | $d = 10$         | $d = 20$      | $d = 50$      | $d = 10$       | $d = 20$      | $d = 50$      |
| $\tau = \frac{1}{4}$ | WA-GAN | <b>0.0103</b>    | <b>0.0074</b> | <b>0.0054</b> | 67.311         | <b>62.905</b> | <b>42.825</b> |
|                      | HTGAN  | 0.0146           | 0.0185        | 0.0199        | 1479.9         | 3106.7        | 1170.7        |
|                      | GPGAN  | 0.0262           | 0.0321        | 0.0326        | <b>66.663</b>  | 65.429        | 52.133        |
| $\tau = \frac{1}{2}$ | WA-GAN | <b>0.0176</b>    | 0.0207        | <b>0.0097</b> | <b>66.734</b>  | <b>52.832</b> | <b>32.468</b> |
|                      | HTGAN  | 0.0247           | <b>0.0179</b> | 0.0126        | 2725.3         | 1447.9        | 2367.4        |
|                      | GPGAN  | 0.0511           | 0.0671        | 0.0672        | 67.841         | 55.308        | 46.996        |
| $\tau = \frac{3}{4}$ | WA-GAN | <b>0.0208</b>    | <b>0.0266</b> | <b>0.0158</b> | <b>59.548</b>  | <b>44.602</b> | <b>31.808</b> |
|                      | HTGAN  | 0.0284           | 0.0338        | 0.055         | 891.62         | 612.44        | 440.45        |
|                      | GPGAN  | 0.0497           | 0.0726        | 0.0818        | 59.919         | 53.758        | 45.541        |

Table 1: Logistic dependence structure. Dependence and extremes scores for different dependence  $\tau$  and dimensions  $d$ . For each scenario and each score, the best score (i.e., the lowest) is indicated in bold.

These results show that WA-GAN is competitive with other methods, from the tail dependence perspective as well as from the quality of the generated extremes. The advantage of WA-GAN seems to increase as the dimension  $d$  increases. Large values of  $d$  is often at the center of interest when generative models for the dependence are used. Note also that it may seem surprising at first that the Wasserstein scores improve as  $d$  increases. This can be explained by several factors. One possible explanation in the present setting is that the dependence structure of the data-generating process is identical across all pairs of variables, which makes it easier to learn as the dimensionality of the space increases. Consequently, to ensure that the comparison remains meaningful, the main interest lies in comparing the values for a given dimension.

Some additional comments can be made for each method. For WA-GAN, the penalization coefficient  $\rho$  in (19) is typically selected to be large among the possible values ( $\rho = 1$  or  $\rho = 3$ ) showing that it is a good idea to enforce marginal constraints (3) in the generative model. For HTGAN, the obtained values of the hyperparameters match the discussion on pages 13–15 in the paper [25] proposing this method. In particular, we get big batch sizes, big latent space dimensions, and low values of the learning rate included in the range recommended by the authors. For GPGAN, the training phase is quicker since, as illustrated on Figure 2, we have a lower training size in this case as this method is trained on points with  $|\mathbf{V}|_\infty > n/k$  while WA-GAN is trained on points with  $|\mathbf{V}|_1 > n/k$ , which is less restrictive. This can be an asset if we have enough data, but can harm GPGAN’s performance if  $d$  is large and  $n$  is not too big.

As an illustration of the performances of the different methods to accurately reproduce

the tail dependence of the data, we plot the empirical bivariate and trivariate extremal coefficients ( $|J| = 2$  and  $|J| = 3$  in (21)) of the test set against the ones of the training set, and against those of all the generative methods for  $d \in \{10, 50\}$  and  $\tau = 0.5$ . In this setting, we know that the true values of these coefficients are  $\sqrt{|J|}$ . This is displayed in Figure 3 for  $d \in \{10, 50\}$ . There is no clear visual difference between WA-GAN and HTGAN, both performing well at representing the coefficients of the test set. GPGAN, however, fails to do so, and even more when the dimension increases. When  $d$  gets larger, the empirical extremal coefficients are below their true values, even in the test set. The reason is that the empirical estimator of  $\Phi$  in (21) becomes less accurate as  $d$  increases [12].

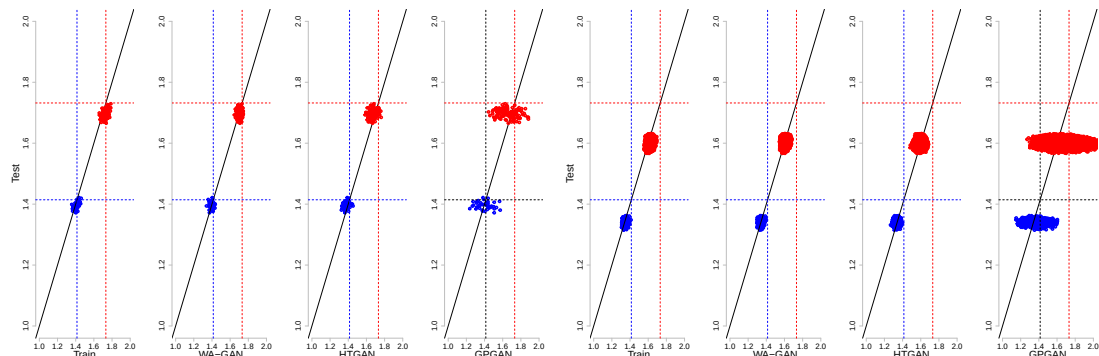


Figure 3: Logistic dependence structure. Left:  $d = 10$ ; right:  $d = 50$ . – Extremal coefficients of order  $k = 2$  (blue) and  $k = 3$  (red). Solid black line is the diagonal. Dashed blue and red lines are the true values of the coefficients for  $k = 2$  and  $k = 3$  respectively.

Next, to illustrate the accuracy of generated extremes by each method, we plot the projection of each generated sample on the first two margins and compare it to the first two margins of the extremes in the test set. We take again  $\tau = 0.5$ . Results for are displayed on Figure 4. WA-GAN seems to perform well, whatever the dimension. The results of HTGAN are less satisfying as it is clear from the figure that the angular measure associated with this method is discrete (this corresponds to the “directions” that we observe in the generated extremes) and this does not match the true angular measure which is logistic. This is theoretically justified by Proposition 2.10 in [25] introducing this method, and already visible from their Figure 2 on page 14. GPGAN also performs well even though its performance decreases in comparison to WA-GAN as the dimension increases; see Table 1.

#### 4.2.2 Hüsler–Reiss dependence structure

In this setup, the random vector  $\mathbf{X}$  is generated using a  $d$ -dimensional multivariate Pareto distribution with Hüsler–Reiss dependence structure having a random parameter matrix and Pareto(2) margins. The parameter matrix generation and the random sampling are done using the R package `graphicalExtremes` [22]. The family is parametrized by a

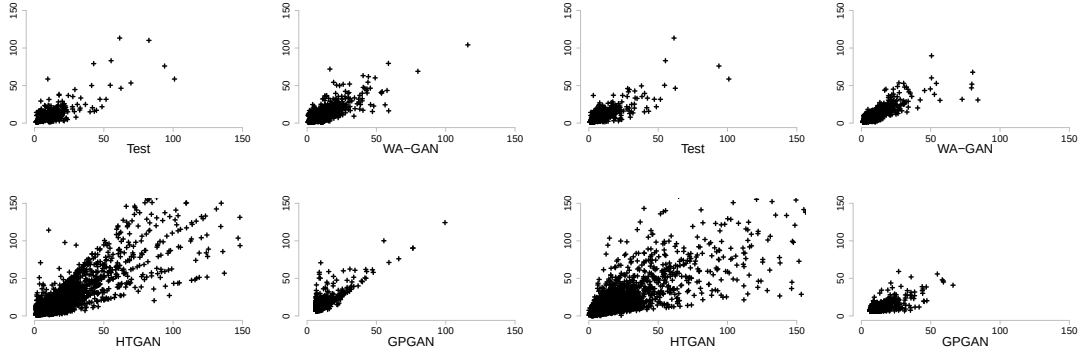


Figure 4: Logistic dependence structure. Left:  $d = 10$ ; right:  $d = 50$ . – First two margins of extremes in the test set and the different generative methods.

variogram matrix  $\Gamma = (\gamma_{ij})_{i,j=1}^d$  where  $\gamma_{ij} = 0$  means complete dependence and  $\gamma_{ij} = +\infty$  means asymptotic independence.

We consider various scenarios in which we assess the performance of WA-GAN and compare it to GPGAN, whose performance in generating multivariate extremes, according to the previous section, is quite close to the results of our procedure. The models are trained on  $n_{\text{train}} = 10\,000$  observations and validated on  $n_{\text{val}} = 5\,000$  observations. The performance assessment is done on  $n_{\text{test}} = 20\,000$  points.

To illustrate the variety of tail strength dependence in the different scenarios, we plot the parameter matrix of each case as heatmaps in Appendix E in the supplement.

**Results.** The setup and the meaning of the scores are exactly the same as explained above Table 1.

| Model  | Dependence score |               |               | Extremes score |              |               |
|--------|------------------|---------------|---------------|----------------|--------------|---------------|
|        | $d = 10$         | $d = 20$      | $d = 50$      | $d = 10$       | $d = 20$     | $d = 50$      |
| WA-GAN | <b>0.0145</b>    | <b>0.0176</b> | <b>0.0123</b> | <b>16.69</b>   | <b>9.512</b> | <b>14.926</b> |
| GPGAN  | 0.0629           | 0.1431        | 0.0989        | 20.979         | 11.406       | 21.067        |

Table 2: Hüsler–Reiss dependence structure. Dependence and extremes scores for different dimensions  $d$ . For each dimension and each score, the best score (i.e., the lowest) is indicated in bold.

WA-GAN remains highly competitive in this more complex scenario, both in terms of dependence and generated extremes. Even though the degrees of dependence vary considerably in this setting (Figure 5), WA-GAN is able to capture them quite accurately, even in high dimensions. For GPGAN, the situation is less favorable, as the deviation from the reference coefficients becomes quite pronounced when  $d$  increases.

We investigate WA-GAN’s sensitivity to the assumption that the angular measure is

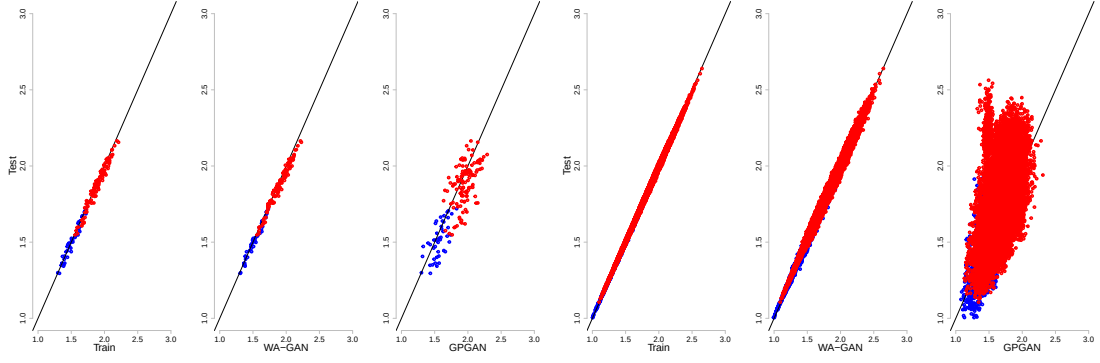


Figure 5: Hüsler–Reiss dependence structure. Left:  $d = 10$ ; right:  $d = 50$ . – Extremal coefficients of order  $k = 2$  (blue) and  $k = 3$  (red).

concentrated on  $\Delta_{d-1}^o$ . Indeed, by construction, the method does not allow the simulation of extremes associated with directions on the boundary of the simplex  $\Delta_{d-1}$ . Based on the parameter matrices shown in Figure 1 in Appendix E, we consider, for  $d \in \{10, 50\}$ , two scenarios exhibiting high and low tail dependence between two components, and we compare the generated extremes to the test extremes in Figure 6. WA-GAN performs well when there is moderate to strong tail dependence but fails to generate extremes that are too close to the axes, corresponding to the low tail dependence scenario, especially if  $d = 50$ . We discuss this issue in the next section under the asymptotic independence scenario and propose a possible improvement to the method.

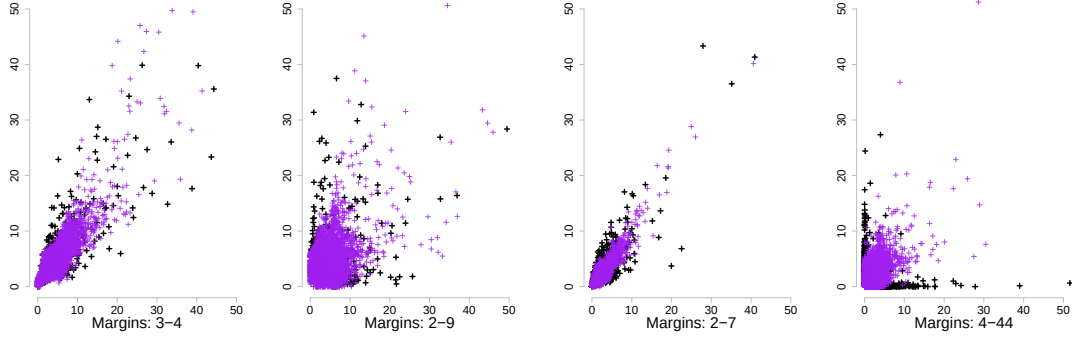


Figure 6: Hüsler–Reiss dependence structure. Left:  $d = 10$ ; right:  $d = 50$ . – Test extremes are in black while generated extremes by WA-GAN are in purple. For each dimension, the left plot corresponds to a high tail dependence scenario while the right plot corresponds to a low tail dependence scenario.

### 4.2.3 Gaussian dependence structure

In this setup, the random vector  $\mathbf{X}$  is generated using a  $d$ -dimensional exchangeable Gaussian copula with correlation parameter  $\rho \in [-1, 1]$  and Pareto distributed marginals with parameter  $\alpha > 0$ . It is easy to verify that this data-generating process does not satisfy our assumptions as Assumption 2 is not verified: all variables are asymptotically independent and the angular measure is concentrated on the vertices of the unit simplex.

We consider various scenarios in which we assess the performance of WA-GAN. We take  $d \in \{10, 20, 50\}$  and a Kendall's tau of  $\tau \in \{1/4, 1/2, 3/4\}$ , corresponding to the correlation parameters  $\rho \in \{\sin(\pi\tau/2) : \tau = 0.25, 0.5, 0.75\}$ . For the margins, we consider  $\alpha = 2$  in any scenario. The models are trained on  $n_{\text{train}} = 10\,000$  observations and validated on  $n_{\text{val}} = 5\,000$  observations. The performance assessment is done on  $n_{\text{test}} = 20\,000$  points. We do not report other methods' performances, as our aim is mainly to see how we may adapt our methodology to be resilient against asymptotic independence in the data.

As our method relies on Aitchison coordinates for the open simplex, it is not able to produce new angles located exactly on the boundary of the simplex. However, one crucial observation is the following: lower tail dependence corresponds to heavier tails in the Aitchison coordinate space. This is illustrated in Figure 7 with  $d = 3$  and  $n = 1000$  data points sampled from a Dirichlet distribution with parameter vector  $\alpha = (2, 2, 2)$  (top) and  $\alpha = (0.1, 0.1, 0.1)$  (bottom). In each situation, we display the raw simplex data, the corresponding coordinates in the orthonormal basis introduced in Proposition 2, and the frequency histogram of the first coordinate values.

Given this observation, we propose to allow for other latent distributions than the standard normal one as input to the generator of WA-GAN. Here, we consider the Student distribution with possible degrees of freedom equal to 1, 2.5, 5, or  $\infty$ , the latter value corresponding to the Gaussian distribution as before. The latent distribution choice is considered as a hyperparameter and is chosen randomly during the random search.

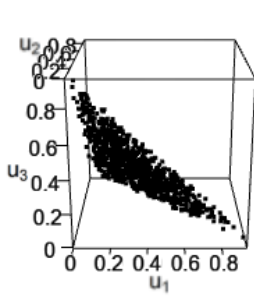
**Results.** Table 3 reports the test scores, for the same two scores as used previously, of WA-GAN in each of the considered scenarios.

| $\tau$               | Dependence score |          |          | Extremes score |          |          |
|----------------------|------------------|----------|----------|----------------|----------|----------|
|                      | $d = 10$         | $d = 20$ | $d = 50$ | $d = 10$       | $d = 20$ | $d = 50$ |
| $\tau = \frac{1}{4}$ | 0.0185           | 0.0214   | 0.0259   | 36.939         | 63.489   | 45.041   |
| $\tau = \frac{1}{2}$ | 0.0335           | 0.0173   | 0.0166   | 53.859         | 117.450  | 71.547   |
| $\tau = \frac{3}{4}$ | 0.0135           | 0.0115   | 0.0108   | 39.144         | 99.651   | 74.948   |

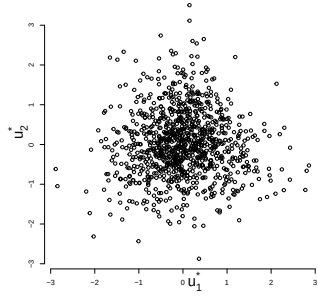
Table 3: Gaussian dependence structure. Dependence and extremes scores for different dependence  $\tau$  and dimensions  $d$ .

The dependence score values are quite similar to those obtained with the logistic model. However, regarding extreme values, the results are worse, as the method tends

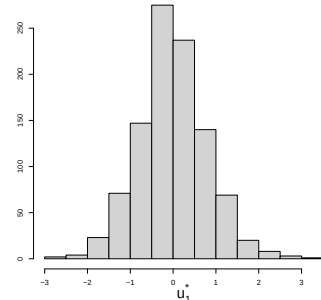




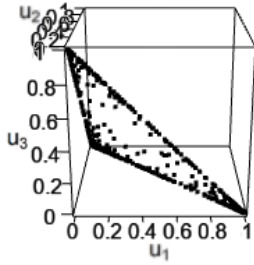
(a) Raw simplex data



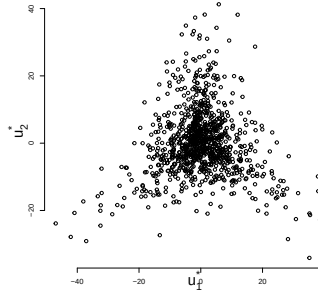
(b) Aitchison coordinates



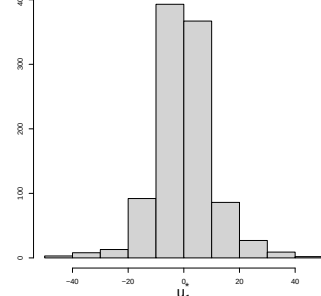
(c) Histogram of the first coordinate values



(d) Raw simplex data



(e) Aitchison coordinates

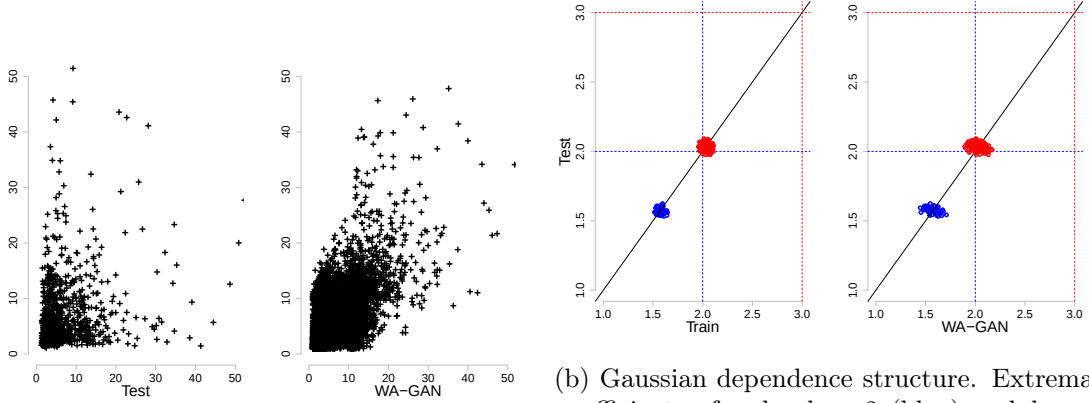


(f) Histogram of the first coordinate values

Figure 7: The Aitchison transform for Dirichlet distributed data in  $d = 3$  with parameter vector  $\alpha = (2, 2, 2)$  (top row)  $\alpha = (0.1, 0.1, 0.1)$  (bottom row).

to produce overly dependent tail observations (Figure 8a). One possible explanation for the fact that the dependence score is not strongly affected is that, under asymptotic independence in the tail, the coefficients are very poorly estimated—even in the test set—so that the score values are not truly reliable in this case (Figure 8b).

Despite the rather poor performance in generating asymptotically independent extremes, it is remarkable that, out of the nine possible scenarios, six of the validated architectures employed a Student rather than Gaussian latent distribution for the Aitchison coordinates. This suggests that allowing for heavier-tailed distributions may be beneficial when asymptotic independence is suspected. The selected degrees of freedom were four times 2.5 and two times 5, indicating that an excessively heavy-tailed distribution, with no finite mean, might be somewhat too extreme.



(a) Gaussian dependence structure. First two margins of extremes in the test set and the generated ones using WA-GAN for  $d = 10$ .

(b) Gaussian dependence structure. Extremal coefficients of order  $k = 2$  (blue) and  $k = 3$  (red). Solid black line is the diagonal. Dashed blue and red lines are the true values of the coefficients for  $k = 2$  and  $k = 3$  respectively.

Figure 8: Results for Gaussian dependence structure with  $d = 10$ .

### 4.3 Financial data: daily returns of industry portfolios

We apply our methodology to analyze the tail behavior of the “value-averaged” daily returns of  $d = 30$  industry portfolios compiled and posted as part of the Kenneth French Data Library. The data in consideration span between 1950 and 2015 with  $n = 16\,694$  observations. This data set is very widely used; for analyses in an extreme value context see e.g. [34] and [13]. Details about the portfolio constructions can be found online.<sup>1</sup> Since we are interested in extreme losses we first multiply all returns by  $-1$ .

As illustrated by the Kendall’s correlation matrix in Figure 9, the daily losses are positively related between the portfolios. We aim to investigate this dependence in the tail part of the distribution.

To sample extremes from the joint distribution, we first compute the maximum likelihood estimators of the univariate GP parameters for each margin. Boxplots illustrating the values of those parameters are reported in Figure 2 in Appendix E. The estimated shape values clearly indicates the presence of heavy tails in the marginal distributions. The GP qq-plots of the marginal fits are to be found on Figures 3, 4 and 5 in Appendix E in the supplement. They indicate a relatively good fit, even though some extreme quantiles are larger than the ones associated with the GP model.

For comparison, we start by computing the test scores for WA-GAN, HTGAN and GPGAN. The training and validation procedures are exactly the same as described in Section 4.2. Here, we train on  $n_{\text{train}} = 7\,000$ , we validate on  $n_{\text{val}} = 3\,000$  and compute the scores on  $n_{\text{test}} = 6\,694$  observations. We again consider  $k = \sqrt{n_{\text{train}}}$  for WA-GAN. For HTGAN, we do not compute the scores on extremes because the marginal distributions

<sup>1</sup>[https://mba.tuck.dartmouth.edu/pages/faculty/ken.french/Data\\_Library/det\\_30\\_ind\\_port.html](https://mba.tuck.dartmouth.edu/pages/faculty/ken.french/Data_Library/det_30_ind_port.html)

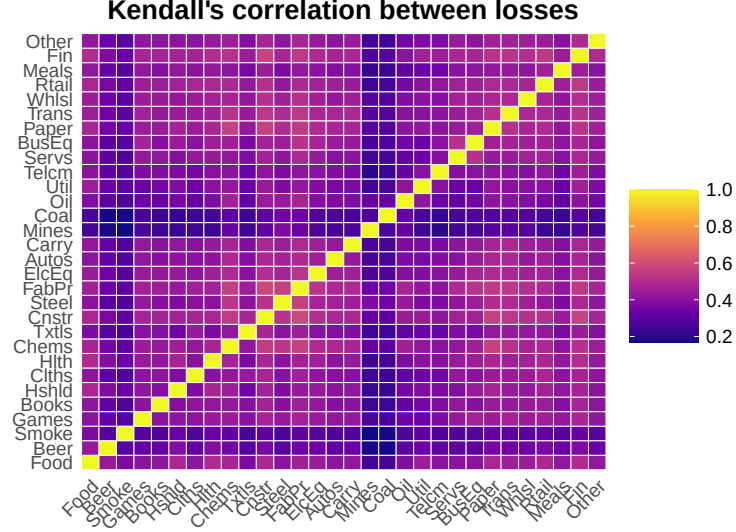


Figure 9: Kendall’s correlation matrix between daily losses of the portfolios.

are not known, see Appendix D. Table 4 shows that WA-GAN seems to be the best at capturing the dependence between extremes in the data while GPGAN has a better 2-Wasserstein distance score.

| method | dependence   | extremes    |
|--------|--------------|-------------|
| WA-GAN | <b>0.018</b> | 11.93       |
| HTGAN  | 0.026        | /           |
| GPGAN  | 0.043        | <b>8.46</b> |

Table 4: Test scores of each method on the  $d = 30$  financial dataset. The best score (i.e., the lowest) is indicated in bold.

As already pointed out in the analysis of [34], the extreme losses of those portfolios can be clustered into different kind of industries. For example, the tobacco industry exhibits asymptotic independence to all other categories (e.g., to the textile industries) while the extreme losses of the business and IT-related industries are dependent. We illustrate the fact that WA-GAN also captures those phenomena in Figure 10a by displaying the two-dimensional projections of the estimated angular measure on those margins. We see that the generated “angles” associated with tobacco/textile portfolios are much more concentrated around the axes than for business/IT, showing that WA-GAN is able to capture quite reasonably this complex tail dependence in this setting.

As a further illustration we consider a portfolio obtained by combining the mines industry portfolio with the coal industry portfolio (with equal weight). It is intuitive and has been shown in [13, 34] that the mines and coal industries exhibit asymptotic

dependence. If one wants to sample the losses associated to the portfolio we propose, this tail dependence should be taken into account. Figure 10b represents the densities of simulated values of the extreme losses of the combined portfolio (i.e., both marginal losses exceeds a large threshold) based on WA-GAN (in black) and based on independent simulations from the marginal GP distributions (in blue). Not taking the tail dependence into account leads to underestimation of the risk associated with this portfolio.

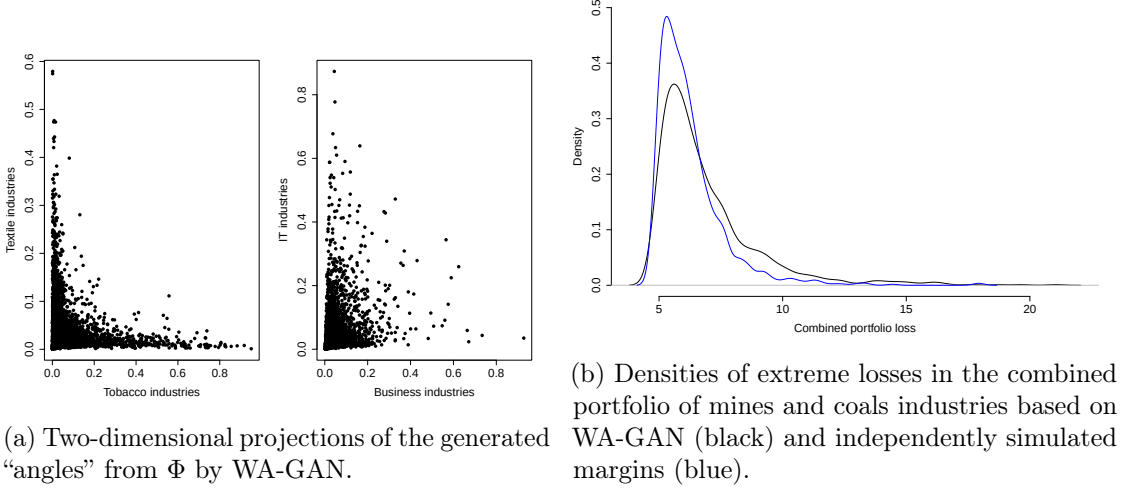


Figure 10: Simulation results for the financial dataset.

## 5 Conclusions

This paper develops the WA-GAN method which is built on the  $L_1$ -norm and on the asymptotic independence of the angular and radial parts of a  $d$ -dimensional regularly varying random vector. An important component is the use of Aitchison coordinates which transforms values on the unit simplex in  $\mathbb{R}^d$  to the entire linear space  $\mathbb{R}^{d-1}$ . The method provides simulated extreme values and can hence be used to estimate probabilities of extreme events, even more extreme than those already observed.

The method is tested on simulated extreme-value datasets of dimensions  $d = 10, 20$ , and  $50$ , with  $10,000$  observations in the training set and  $5,000$  in the validation set. Its performance is evaluated using two metrics: a dependence score based on extremal coefficients and a Wasserstein score illustrating the quality of the generated extremes. In both cases, WA-GAN performs reasonably well compared to other existing methods such as HTGAN and GPGAN. However, its performance may deteriorate under asymptotic independence, as it relies on the Aitchison space, where low tail dependence corresponds to heavier-tailed coordinates, while WA-GAN uses a Gaussian latent distribution. Possible improvements to address this limitation will be discussed in future work.

The methods are also applied to a financial data set from the Kenneth French Data Library. Again WA-GAN performs well and underlines the importance of being able

to handle dependence between different financial portfolios. In terms of performance measures, WA-GAN has the best dependence score, while GPGAN has a better 2-Wasserstein score than WA-GAN.

In this paper we do not quantify the uncertainty in the estimated probability of an extreme event. However, provided enough computational power is available, this could be done using a approach similar to a parametric bootstrap. For this one would repeat the following, say, 100 times: first use the fitted WA-GAN to simulate a new extreme data set of the same size as the original one, then fit a new WA-GAN to this simulated data set and use the new WA-GAN to estimate the probability of interest. This will provide 100 estimated probability values which can be used to find “confidence bounds” for the original estimated probability.

An alternative approach would be to take a Bayesian perspective: choose a prior distribution for the marginal tail parameters and the angular measure (or any other means of describing the dependence structure) and then simulate from the predictive distribution given the data: first simulate random tail and parameters from the posterior distribution and then sample from the resulting multivariate generalized Pareto distribution. The advantage of such an approach is that the estimation uncertainty would be taken into account through the posterior distribution of the tail and dependence parameters. Drawbacks are that the construction of a suitable prior on the dependence structure may become quite delicate and the generation from the posterior distribution rather involved, especially in high dimensions.

Another important aspect, not addressed in this work but potentially relevant in applications, is the possible time-evolving nature of the extremal dependence structure of the data [14, 17], or its variation with respect to a latent variable. Covariates can be incorporated into (Wasserstein) GANs by using their conditional versions [40, 58], which could facilitate the generation of new multivariate extremes for specific covariate values, like time, in a manner similar to what has been explored for image generation in [24]. These directions are left for future work.

## Code and data availability

The code is publicly available on the repository <https://github.com/stephanelh98/extreme-WGAN>. The data underlying the analysis in Section 4.3 is publicly available from the Kenneth R. French data library [https://mba.tuck.dartmouth.edu/pages/faculty/ken.french/data\\_library.html](https://mba.tuck.dartmouth.edu/pages/faculty/ken.french/data_library.html).

## Acknowledgments

The authors sincerely thank the anonymous reviewers for their insightful comments and suggestions, which have significantly improved the quality and clarity of this paper.

The research of Stéphane Lhaut was supported by the Fonds National de la Recherche Scientifique (FNRS, Belgium) within the framework of a FRIA grant (No. 1.E.114.23F). Computational resources have been provided by the Consortium des Équipements de

Calcul Intensif (CÉCI), funded by the Fonds de la Recherche Scientifique de Belgique (F.R.S.-FNRS) under Grant No. 2.5020.11 and by the Walloon Region.

## References

- [1] John Aitchison. The statistical analysis of compositional data. *Journal of the Royal Statistical Society. Series B*, 44(2):139–177, 1982.
- [2] Michaël Allouche, Stéphane Girard, and Emmanuel Gobet. EV-GAN: Simulation of Extreme Events with ReLU Neural Networks. *Journal of Machine Learning Research*, 23(150):1–39, 2022.
- [3] Michaël Allouche, Stéphane Girard, and Emmanuel Gobet. ExcessGAN: simulation above extreme thresholds using Generative Adversarial Networks. *HAL preprint, hal-05044516*, pages 1–24, 2025.
- [4] Hugues Annoye and Cédric Heuchenne. Generating survey databases with Wasserstein Generative Adversarial Networks. *Applied Intelligence*, 55(1095):1–17, 2025.
- [5] Martin Arjovsky, Soumith Chintala, and Léon Bottou. Wasserstein generative adversarial networks. In Doina Precup and Yee Whye Teh, editors, *Proceedings of the 34th International Conference on Machine Learning*, volume 70 of *Proceedings of Machine Learning Research*, pages 214–223. PMLR, 06–11 Aug 2017.
- [6] August A. Balkema and Laurens de Haan. Residual life time at great age. *The Annals of Probability*, 2(5):792–804, 1974.
- [7] Jan Beirlant, Yuri Goegebeur, Johan Segers, and Jozef L. Teugels. *Statistics of Extremes: Theory and Applications*. Wiley Hoboken, 2004.
- [8] Siddharth Bhatia, Arjit Jain, and Bryan Hooi. ExGAN: Adversarial generation of extreme samples. In *Proceedings of the The Thirty-Fifth AAAI Conference on Artificial Intelligence*, volume 35, pages 6750–6758, 2021.
- [9] Gérard Biau, Maxime Sangnier, and Ugo Tanielian. Some theoretical insights into Wasserstein GANs. *J. Mach. Learn. Res.*, 22:45, 2021. Id/No 119.
- [10] Christopher M. Bishop. *Deep Learning*. Springer, 2024.
- [11] Younes Boulaguiem, Jakob Zscheischler, Edoardo Vignotto, Karin van der Wiel, and Sebastian Engelke. Modeling and simulating spatial extremes by combining extreme value theory with generative adversarial networks. *Environmental Data Science*, 1:1–18, 2022.
- [12] Stéphan Cléménçon, Hamid Jalalzai, Stéphane Lhaut, Anne Sabourin, and Johan Segers. Concentration bounds for the empirical angular measure with statistical learning applications. *Bernoulli*, 29(4):2797–2827, 2023.

- [13] Daniel Cooley and Emeric Thibaud. Decompositions of dependence for high-dimensional extremes. *Biometrika*, 106(3):587–604, 2019.
- [14] Miguel de Carvalho and Karla Vianey Palacios Ramirez. Semiparametric Bayesian modelling of nonstationary joint extremes: How do big tech’s extreme losses behave? *Journal of the Royal Statistical Society Series C: Applied Statistics*, 74:447–465, 2025.
- [15] Laurens de Haan and Ana Ferreira. *Extreme Value Theory: An Introduction*. Springer New York, 2006.
- [16] Laurens de Haan and Sidney I. Resnick. Limit theory for multivariate sample extremes. *Zeitschrift für Wahrscheinlichkeitstheorie und verwandte Gebiete*, 40:317–337, 1977.
- [17] Holger Drees. Statistical inference on a changing extreme value dependence structure. *The Annals of Statistics*, 51(4):1824–1849, 2023.
- [18] Juan J. Egozcue, Vera Pawlowsky-Glahn, Gloria Mateu-Figueras, and Carles Barceló-Vidal. Isometric logratio transformations for compositional data analysis. *Mathematical Geology*, 35(3):279–300, 2003.
- [19] John H.J. Einmahl, Laurens de Haan, and Vladimir I. Piterbarg. Nonparametric estimation of the spectral measure of an extreme value distribution. *Annals of Statistics*, 29(5):1401–1423, 2001.
- [20] John H.J. Einmahl, Laurens de Haan, and Ashoke Kumar Sinha. Estimating the spectral measure of an extreme value distribution. *Stochastic Processes and their Applications*, 70(2):143–171, 1997.
- [21] John H.J. Einmahl and Johan Segers. Maximum empirical likelihood estimation of the spectral measure of an extreme-value distribution. *Annals of Statistics*, 37(5B):2953–2989, 2009.
- [22] Sebastian Engelke, Adrien S. Hitz, Nicola Gnecco, and Manuel Hentschel. *graphicalExtremes: Statistical Methodology for Graphical Extreme Value Models*, 2024. R package version 0.3.2, <https://CRAN.R-project.org/package=graphicalExtremes>.
- [23] Ana Ferreira and Laurens de Haan. The generalized Pareto process; with a view towards application and simulation. *Bernoulli*, 20(4):1717 – 1737, 2014.
- [24] Jon Gauthier. Conditional generative adversarial nets for convolutional face generation. Technical report, Stanford University, 2015.
- [25] Stéphane Girard, Emmanuel Gobet, and Jean Pachebat. Deep generative modeling of multivariate dependent extremes. *HAL preprint*, *hal-04700084v2*, pages 1–35, 2024.

- [26] Ian Goodfellow, Jean Pouget-Abadie, Mehdi Mirza, Bing Xu, David Warde-Farley, Sherjil Ozair, Aaron Courville, and Yoshua Bengio. Generative adversarial nets. In Z. Ghahramani, M. Welling, C. Cortes, N. Lawrence, and K.Q. Weinberger, editors, *Advances in Neural Information Processing Systems*, volume 27. Curran Associates, Inc., 2014.
- [27] Ishaan Gulrajani, Faruk Ahmed, Martin Arjovsky, Vincent Dumoulin, and Aaron C Courville. Improved training of Wasserstein GANs. In I. Guyon, U. Von Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc., 2017.
- [28] Ali Hasan, Khalil Elkhail, Yuting Ng, João M. Pereira, Sina Farsiu, Jose H. Blanchet, and Vahid Tarokh. Modeling extremes with  $d$ -max-decreasing neural networks. In *Proceedings of the 38th Conference on Uncertainty in Artificial Intelligence (UAI)*, volume 180 of *Proceedings of Machine Learning Research*, pages 759–768. PMLR, 2022.
- [29] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Delving deep into rectifiers: Surpassing human-level performance on imagenet classification. In *Proceedings of the IEEE International Conference on Computer Vision (ICCV)*, pages 1026–1034, 2015.
- [30] Martin Heusel, Hubert Ramsauer, Thomas Unterthiner, Bernhard Nessler, and Sepp Hochreiter. GANs trained by a two time-scale update rule converge to a local Nash equilibrium. In I. Guyon, U. Von Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc., 2017.
- [31] Tennessee Hickling and Dennis Prangle. Flexible Tails for Normalizing Flows. *arXiv:2406.16971v2*, 2025.
- [32] Chenglei Hu and Daniela Castro-Camilo. GPDFlow: Generative Multivariate Threshold Exceedance Modeling via Normalizing Flows. *arXiv:2503.11822v1*, 2025.
- [33] Todd Huster, Jeremy Cohen, Zinan Lin, Kevin Chan, Charles Kamhoua, Nandi O. Leslie, Cho-Yu Jason Chiang, and Vyas Sekar. Pareto GAN: Extending the representational power of GANs to heavy-tailed distributions. In Marina Meila and Tong Zhang, editors, *Proceedings of the 38th International Conference on Machine Learning*, volume 139 of *Proceedings of Machine Learning Research*, pages 4523–4532, 2021.
- [34] Anja Janßen and Phyllis Wan.  $k$ -means clustering of extremes. *Electronic Journal of Statistics*, 14(1):1211–1233, 2020.
- [35] Diederik Kingma and Jimmy Ba. Adam: A method for stochastic optimization. In *International Conference on Learning Representations*, 2015.
- [36] Nicolas Lafon, Philippe Naveau, and Ronan Fablet. A VAE approach to sample multivariate extremes. *arXiv:2306.10987*, 2023.



- [37] Jiawei Li, Dong Li, Peng Li, and Gennady Samorodnitsky. Generalized Pareto GAN: Generating extremes of distributions. In *Proceedings of the 2024 International Joint Conference on Neural Networks (IJCNN)*, pages 1–8. IEEE, 2024.
- [38] Andrew L. Maas, Awni Y. Hannun, and Andrew Y. Ng. Rectifier nonlinearities improve neural network acoustic models. In *Proceedings of the 30th International Conference on Machine Learning (ICML)*, 2013. Workshop on Deep Learning for Audio, Speech and Language Processing.
- [39] Thomas Mikosch and Olivier Wintenberger. *Extreme Value Theory for Time Series*. Springer Series in Operations Research and Financial Engineering. Springer, 2024.
- [40] Mehdi Mirza and Simon Osindero. Conditional generative adversarial nets. *arxiv.org:1411.1784*, 2014.
- [41] Lambert De Monte, Raphaël Huser, Ioannis Papastathopoulos, and Jordan Richards. Generative modelling of multivariate geometric extremes using normalising flows. *arXiv:2505.02957v1*, 2025.
- [42] Anas Mourahib, Anna Kiriliouk, and Johan Segers. Multivariate generalized Pareto distributions along extreme directions. *Extremes*, 28:239–272, 2025.
- [43] Philippe Naveau and Johan Segers. Multivariate extreme value theory. *arXiv.2412.18477*, 2024.
- [44] Adam Paszke, Sam Gross, Francisco Massa, Adam Lerer, James Bradbury, Gregory Chanan, Trevor Killeen, Zeming Lin, Natalia Gimelshein, Luca Antiga, Alban Desmaison, Andreas Kopf, Edward Yang, Zachary DeVito, Martin Raison, Alykhan Tejani, Sasank Chilamkurthy, Benoit Steiner, Lu Fang, Junjie Bai, and Soumith Chintala. Pytorch: An imperative style, high-performance deep learning library. In H. Wallach, H. Larochelle, A. Beygelzimer, F. d'Alché-Buc, E. Fox, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019.
- [45] James Pickands III. Statistical inference using extreme order statistics. *The Annals of Statistics*, 3(1):119–131, 1975.
- [46] R Core Team. *R: A Language and Environment for Statistical Computing*. R Foundation for Statistical Computing, Vienna, Austria, 2023.
- [47] Sidney I. Resnick. *Extreme Values, Regular Variation and Point Processes*. Springer New York, 1987.
- [48] Sidney I. Resnick. *Heavy-Tail Phenomena: Probabilistic and Statistical Modeling*. Springer New York, 2007.
- [49] Sidney I. Resnick. *The Art of Finding Hidden Risks*. Springer, 2024.

- [50] Holger Rootzén, Johan Segers, and Jennifer L. Wadsworth. Multivariate generalized Pareto distributions: Parametrizations, representations, and properties. *Journal of Multivariate Analysis*, 165:117–131, 2018.
- [51] Holger Rootzén, Johan Segers, and Jennifer L. Wadsworth. Multivariate peaks over thresholds models. *Extremes*, 21(1):115–145, 2018.
- [52] Holger Rootzén and Nader Tajvidi. Multivariate generalized Pareto distributions. *Bernoulli*, 12(5):917–930, 2006.
- [53] Dominic Schuhmacher, Björn Bähre, Nicolas Bonneel, Carsten Gottschlich, Valentin Hartmann, Florian Heinemann, Bernhard Schmitzer, and Jörn Schrieber. *transport: Computation of Optimal Transport Plans and Wasserstein Distances*, 2024. R package version 0.15-4.
- [54] Richard L Smith. Estimating tails of probability distributions. *The Annals of Statistics*, 15(3):1174–1207, 1987.
- [55] Aad W. van der Vaart and Jon A. Wellner. *Weak Convergence and Empirical Processes: With Applications to Statistics*. Springer, New York, 1996.
- [56] Cédric Villani. *Topics in Optimal Transportation*, volume 58 of *Graduate Studies in Mathematics*. American Mathematical Society, 2003.
- [57] Jakob Benjamin Wessel, Callum J. R. Murphy-Barltrop, and Emma S. Simpson. A comparison of generative deep learning methods for multivariate angular simulation. *arxiv.org:2504.21505*, 2025.
- [58] Ming Zheng, Tong Li, Rui Zhu, Yahui Tang, Mingjing Tang, Leilei Lin, and Zifei Ma. Conditional Wasserstein generative adversarial network-gradient penalty-based approach to alleviating imbalanced data classification. *Information Sciences*, 512:1009–1023, 2020.

## A Proof of Proposition 1

We start by noting that  $\mathbf{X} \stackrel{\mathcal{L}}{=} \mathbf{b}(\mathbf{V})$  and, by the continuity of the margins  $F_j$ , each tail quantile function  $b_j$  is strictly increasing, so that  $\{\mathbf{X} \not\leq \mathbf{b}(t)\} = \{\mathbf{V} \not\leq t\}$  for any  $t > 1$ . Consequently,

$$\frac{\mathbf{X} - \mathbf{b}(t)}{\mathbf{a}(t)} \Big| \mathbf{X} \not\leq \mathbf{b}(t) \stackrel{\mathcal{L}}{=} \frac{\mathbf{b}(t(t^{-1}\mathbf{V})) - \mathbf{b}(t)}{\mathbf{a}(t)} \Big| \mathbf{V} \not\leq t.$$

Furthermore, it follows easily from Assumption 1 that

$$(t^{-1}\mathbf{V} \mid \mathbf{V} \not\leq t) \xrightarrow{\mathcal{L}} \mathbf{Y}, \quad t \rightarrow \infty, \tag{25}$$

where, for Borel set  $B \subseteq \mathbb{E}$ , writing  $\mathbb{L}_1 \stackrel{\text{def}}{=} \{\mathbf{x} \in [0, \infty)^d : \mathbf{x} \not\leq 1\}$ , we have

$$\mathbf{P}(\mathbf{Y} \in B) = \frac{\nu(B \cap \mathbb{L}_1)}{\nu(\mathbb{L}_1)}. \quad (26)$$

The exponent measure  $\nu$  is connected to the angular measure  $\Phi$  introduced in (2). More precisely [7, Chapter 8], for any bounded, measurable function  $f : \mathbb{E} \rightarrow \mathbb{R}$  that is zero in a neighborhood of the origin, we have

$$\nu(f) \stackrel{\text{def}}{=} \int_{\mathbb{E}} f(x) d\nu(x) = d \int_{\Delta_{d-1}} \int_0^\infty f(y\mathbf{w}) \frac{dy}{y^2} d\Phi(\mathbf{w}). \quad (27)$$

In particular, for any for Borel set  $B \subseteq \mathbb{E}$ ,

$$\begin{aligned} d^{-1}\nu(B \cap \mathbb{L}_1) &= \int_{\Delta_{d-1}} \int_0^\infty \mathbb{I}_{B \cap \mathbb{L}_1}(y\mathbf{w}) \frac{dy}{y^2} d\Phi(\mathbf{w}) \\ &= \int_{\Delta_{d-1}} \int_1^\infty \mathbb{I}_{B \cap \mathbb{L}_1}(y\mathbf{w}) \frac{dy}{y^2} d\Phi(\mathbf{w}) \\ &= \mathbf{P}(Y\boldsymbol{\Theta} \in B \cap \mathbb{L}_1), \end{aligned}$$

where  $Y$  is unit-Pareto distributed and independent of  $\boldsymbol{\Theta} \sim \Phi$ . Hence, by (26),

$$\mathbf{P}(\mathbf{Y} \in B) = \frac{d \mathbf{P}(Y\boldsymbol{\Theta} \in B \cap \mathbb{L}_1)}{d \mathbf{P}(Y\boldsymbol{\Theta} \in \mathbb{L}_1)} = \mathbf{P}(Y\boldsymbol{\Theta} \in B \mid Y\boldsymbol{\Theta} \in \mathbb{L}_1),$$

From the local uniformity of the convergence in point (ii) of Remark 1, guaranteed by Assumption 3, and the fact that  $\mathbf{P}(Y_j > 0) = 1$  for any  $j \in \{1, \dots, d\}$  by the previous computations and Assumption 2, the extended continuous mapping theorem [55, Theorem 1.11.1] applies, yielding

$$\frac{\mathbf{b}(t(t^{-1}\mathbf{V})) - \mathbf{b}(t)}{\mathbf{a}(t)} \Big| \mathbf{V} \not\leq t \xrightarrow{\mathcal{L}} \frac{\mathbf{Y}^\xi - 1}{\boldsymbol{\xi}}, \quad t \rightarrow \infty.$$

This concludes the proof of (7) and (8). Eq. (9) is a direct consequence of (7) as

$$\sigma_j(u_j(t)) = a_j \left( \frac{1}{1 - F_j(b_j(t))} \right) = a_j(t), \quad t > 1$$

since  $F_j(F_j^{-1}(p)) = p$  for all  $0 < p < 1$  by continuity of  $F_j$ .  $\square$

## B The Aitchison simplex

The open simplex  $\Delta_{d-1}^\circ$  is equipped with an inner product space structure with the following operations [1]: for  $\mathbf{v} = (v_1, \dots, v_d), \mathbf{w} = (w_1, \dots, w_d) \in \Delta_{d-1}^\circ$  and  $\alpha \in \mathbb{R}$ , define

the following operations:

$$\begin{aligned}
\mathbf{v} \oplus \mathbf{w} &\stackrel{\text{def}}{=} \left( \frac{v_j w_j}{\sum_{i=1}^d v_i w_i} \right)_{j=1}^d && \text{(addition),} \\
\alpha \odot \mathbf{v} &\stackrel{\text{def}}{=} \left( \frac{v_j^\alpha}{\sum_{i=1}^d v_i^\alpha} \right)_{j=1}^d && \text{(scalar multiplication),} \\
\langle \mathbf{v}, \mathbf{w} \rangle_A &\stackrel{\text{def}}{=} \sum_{i=1}^d \log \left( \frac{v_i}{g(\mathbf{v})} \right) \log \left( \frac{w_i}{g(\mathbf{w})} \right) && \text{(Aitchison inner product),}
\end{aligned}$$

where  $g : \mathbf{u} \in \Delta_{d-1}^\circ \mapsto g(\mathbf{u}) \stackrel{\text{def}}{=} (\prod_{i=1}^d u_i)^{1/d}$  is the geometric mean. It is an easy exercise to verify that  $(\Delta_{d-1}^\circ, \oplus, \odot)$  forms a real vector space with zero element  $\mathbf{0}_A = \mathbf{1}_d/d \in \Delta_{d-1}^\circ$ , where  $\mathbf{1}_d = (1, \dots, 1) \in \mathbb{R}^d$ , and that  $\langle \cdot, \cdot \rangle_A$  defines an inner product on it. The open simplex equipped with this structure is often called the *Aitchison simplex*.

A trivial but useful observation is that if one introduces the so-called *Centered LogRatio (CLR)* operator

$$\text{clr} : \mathbf{u} \in \Delta_{d-1}^\circ \mapsto \text{clr}(\mathbf{u}) \stackrel{\text{def}}{=} \left( \log \left( \frac{u_j}{g(\mathbf{u})} \right) \right)_{j=1}^d \in \mathbb{H} \subset \mathbb{R}^d, \quad (28)$$

where  $\mathbb{H} \stackrel{\text{def}}{=} \{\mathbf{x} \in \mathbb{R}^d : \langle \mathbf{x}, \mathbf{1}_d \rangle = 0\}$  with  $\langle \cdot, \cdot \rangle$  the standard inner product in  $\mathbb{R}^d$ , then

$$\langle \mathbf{v}, \mathbf{w} \rangle_A = \langle \text{clr}(\mathbf{v}), \text{clr}(\mathbf{w}) \rangle, \quad \mathbf{v}, \mathbf{w} \in \Delta_{d-1},$$

so that  $\text{clr}$  naturally defines an isometry between  $(\Delta_{d-1}^\circ, \langle \cdot, \cdot \rangle_A)$  and  $(\mathbb{H}, \langle \cdot, \cdot \rangle)$ . The inverse transformation  $\text{clr}^{-1} : \mathbb{H} \mapsto \Delta_{d-1}^\circ$  is the well known *softmax* function

$$\text{clr}^{-1}(\mathbf{x}) = \left( \frac{\exp(\mathbf{x}_j)}{\sum_{i=1}^d \exp(\mathbf{x}_i)} \right)_{j=1}^d = \text{softmax}(\mathbf{x}), \quad \mathbf{x} \in \mathbb{H}.$$

The above considerations lead to a clear and easy way to construct an orthonormal basis of the Aitchison simplex:

1. Take any linearly independent family  $\{\mathbf{x}_1, \dots, \mathbf{x}_{d-1}\}$  in  $\mathbb{H}$ .
2. Apply the Gram–Schmidt procedure to produce an orthonormal basis  $\{\mathbf{e}_1, \dots, \mathbf{e}_{d-1}\}$  of  $\mathbb{H}$  with respect to the standard inner product on  $\mathbb{R}^d$ .
3. Compute  $\{\mathbf{e}_i^* \stackrel{\text{def}}{=} \text{clr}^{-1}(\mathbf{e}_i) : i = 1, \dots, d-1\}$ , which forms an orthonormal basis of the Aitchison simplex.

Applying the above procedure to the free family

$$\{\mathbf{x}_i \stackrel{\text{def}}{=} (0, \dots, 0, 1, -1, 0, \dots, 0) \in \mathbb{H} : i = 1, \dots, d-1\},$$

where in  $\mathbf{x}_i$  the 1 element lies at the  $i$ -th position, leads to the orthonormal basis of the Aitchison simplex presented in Proposition 2.

## C Angular measure with respect to another norm

The angular measure  $\Phi$  in (2) can be introduced with respect to an arbitrary norm  $\|\cdot\|$  on  $\mathbb{R}^d$ , that is, if we let

$$\tilde{R} = \|\mathbf{V}\| \quad \text{and} \quad \tilde{\mathbf{W}} = \tilde{R}^{-1}\mathbf{V},$$

one can consider the measure  $\tilde{\Phi}$  on  $\tilde{\Delta}_{d-1} \stackrel{\text{def}}{=} \{\mathbf{x} \in [0, 1]^d : \|\mathbf{x}\| = 1\}$  obtained by taking, as in (2),

$$\tilde{\Phi}(A) = \lim_{t \rightarrow \infty} \mathbf{P} \left( \tilde{\mathbf{W}} \in A \mid \tilde{R} \geq t \right), \quad (29)$$

for any Borel set  $A \subseteq \tilde{\Delta}_{d-1}$  such that  $\tilde{\Phi}(\partial A) = 0$ . Popular choices for  $\|\cdot\|$  consist of the  $L_p$  norms  $|\cdot|_p$  on  $\mathbb{R}^d$  for  $p \in [1, \infty]$ , see, e.g., [19, 21, 12]. Even though we only consider the norm  $|\cdot|_1$  in Algorithm 1, we show that a simple post-processing step permits to provide an estimate for  $\tilde{\Phi}$  also.

From (29), one may show that  $\tilde{\Phi}(A) = \Psi(A)/\Psi(\tilde{\Delta}_{d-1})$  where  $\Psi$  is the finite Borel measure on  $\tilde{\Delta}_{d-1}$  obtained by taking

$$\Psi(A) = \lim_{t \rightarrow \infty} t \mathbf{P} \left( \tilde{\mathbf{W}} \in A, \tilde{R} \geq t \right)$$

provided  $\Psi(\partial A) = 0$ . From (27), for any bounded, measurable  $f : \mathbb{E} \rightarrow \mathbb{R}$  vanishing in a neighborhood of  $\mathbf{0}$ , we get

$$d \int_{\Delta_{d-1}} \int_0^\infty f(r\mathbf{w}) \frac{dr}{r^2} d\Phi(\mathbf{w}) = \int_{\tilde{\Delta}_{d-1}} \int_0^\infty f(r\tilde{\mathbf{w}}) \frac{dr}{r^2} d\Psi(\tilde{\mathbf{w}}),$$

Taking  $f(x) = g(x/\|x\|)\mathbb{I}(\|x\| \geq 1)$  for  $x \in \mathbb{E}$ , where  $g : \tilde{\Delta}_{d-1} \rightarrow \mathbb{R}$  is some bounded function, we get

$$d \int_{\Delta_{d-1}} g\left(\frac{\mathbf{w}}{\|\mathbf{w}\|}\right) \|\mathbf{w}\| d\Phi(\mathbf{w}) = \int_{\tilde{\Delta}_{d-1}} g(\tilde{\mathbf{w}}) d\Psi(\tilde{\mathbf{w}}).$$

Picking  $g(\tilde{\mathbf{w}}) = \mathbb{I}(\tilde{\mathbf{w}} \in \tilde{\Delta}_{d-1})$  gives

$$\Psi(\tilde{\Delta}_{d-1}) = d \int_{\Delta_{d-1}} \|\mathbf{w}\| d\Phi(\mathbf{w}),$$

so that for any bounded  $g : \tilde{\Delta}_{d-1} \rightarrow \mathbb{R}$  we have

$$\int_{\tilde{\Delta}_{d-1}} g(\tilde{\mathbf{w}}) d\tilde{\Phi}(\tilde{\mathbf{w}}) = \frac{\int_{\Delta_{d-1}} g\left(\frac{\mathbf{w}}{\|\mathbf{w}\|}\right) \|\mathbf{w}\| d\Phi(\mathbf{w})}{\int_{\Delta_{d-1}} \|\mathbf{w}\| d\Phi(\mathbf{w})}. \quad (30)$$

In particular,  $\tilde{\Phi}$  is fully determined by  $\Phi$ .

Suppose that we observe  $\Theta_1, \dots, \Theta_n \sim \Phi$  as it would be (approximately) the case via the trained generator in Algorithm 1. They lead to the empirical estimator  $\hat{\Phi} = (1/n) \sum_{i=1}^n \delta_{\Theta_i}$  for  $\Phi$ . Relation (30) then justifies the approximation

$$\begin{aligned} \int_{\tilde{\Delta}_{d-1}} g(\tilde{\mathbf{w}}) d\tilde{\Phi}(\tilde{\mathbf{w}}) &\approx \frac{\int_{\Delta_{d-1}} g\left(\frac{\mathbf{w}}{\|\mathbf{w}\|}\right) \|\mathbf{w}\| d\hat{\Phi}(\mathbf{w})}{\int_{\Delta_{d-1}} \|\mathbf{w}\| d\hat{\Phi}(\mathbf{w})} \\ &= \sum_{i=1}^n \Lambda_i g\left(\frac{\Theta_i}{\|\Theta_i\|}\right) = \int_{\tilde{\Delta}_{d-1}} g(\tilde{\mathbf{w}}) d\hat{\tilde{\Phi}}(\tilde{\mathbf{w}}), \end{aligned}$$

where

$$\hat{\tilde{\Phi}} \stackrel{\text{def}}{=} \sum_{i=1}^n \Lambda_i \delta_{\Theta_i / \|\Theta_i\|}, \quad \Lambda_i \stackrel{\text{def}}{=} \frac{\|\Theta_i\|}{\sum_{i=1}^n \|\Theta_i\|}.$$

Said otherwise, a sample  $\Theta_1, \dots, \Theta_n \sim \Phi$  can be used to produce an estimator for  $\tilde{\Phi}$  by considering the weighted empirical distribution of the rescaled angles  $\Theta_i$  on  $\tilde{\Delta}_{d-1}$ , where the weights are proportional to the norm  $\|\Theta_i\|$ .