

On the Simulation of Extreme Events with Neural Networks

Generative AI Modelling for Extreme Events

Vianey Palacios



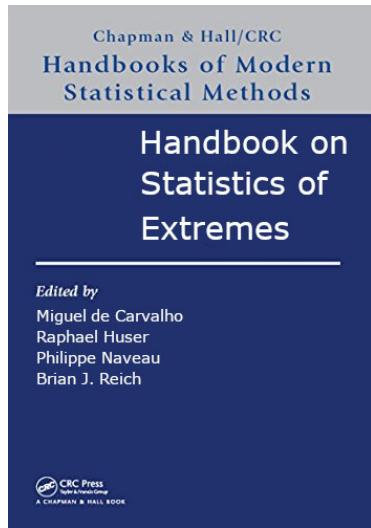
Based on Chapter from Handbook on Statistics of Extremes

This talk is based on:

Chapter 20 – “*On the Simulation of Extreme Events with Neural Networks*”

by Michaël Allouche, Stéphane Girard & Emmanuel Gobet

In *Handbook on Statistics of Extremes*



Introduction, Motivation & Background

Generative Modeling

- This chapter explores AI in the context of extreme events, with a focus on **generative models**.
- It highlights the **limitations** of mainstream AI approaches in simulating **heavy-tailed** phenomena and presents fresh **solutions**.
- This talk focus on **Generative Adversarial Networks (GANs)**.
- The chapter also discussed **Variational Auto-Encoders** and **Difussion Models**.

What is Generative Modeling?

- Given $\{\mathbf{x}_i\}_{i=1}^n$ from an **unknown** distribution $p_{\mathbf{X}}$.
- The **goal** of generative modeling is to find:
 - **latent distribution**: $p_{\mathbf{Z}}$ defined on some set $\mathcal{Z}$.
 - **generator**: $G : \mathcal{Z} \rightarrow \mathcal{X}$.

such that

$$G(\mathbf{Z}) \stackrel{d}{=} \mathbf{X}, \quad \mathbf{Z} \sim p_{\mathbf{Z}}$$

- We **do not** want to **infer** the unknown distribution $p_{\mathbf{X}}$ from the data.

Theorem

(Kuratowski). Let $(\mathcal{Z}, \mu_{\mathbf{Z}})$ and $(\mathcal{X}, \mu_{\mathbf{X}})$ be two Polish probability spaces. Then, there exists a (non-unique) measurable bijection $G : \mathcal{Z} \rightarrow \mathcal{X}$ such that $\mu_{\mathbf{Z}}\{G^{-1}(E)\} = \mu_{\mathbf{X}}(E)$ and $\mu_{\mathbf{X}}\{G(F)\} = \mu_{\mathbf{Z}}(F)$, for all Borel sets $E \subset \mathcal{X}$ and $F \subset \mathcal{Z}$.

- In practice, distributions **easy to simulate** are considered for $p_{\mathbf{Z}}$ (Gaussian, uniform, etc). The **modelling effort** is usually put on the generator.

Generative Modeling Framework

- Focus on a **parametric** family of generators $\{G_{\theta}\}_{\theta \in \Theta}$ where $\Theta \subset \mathbb{R}^p$ such that $G_{\theta}(\mathbf{Z}) \sim P_{\theta}$.
- The problem comes in finding the **'best' parameter** θ^* such that p_{θ} , and p_X are as close as possible:

$$G_{\theta^*}(\mathbf{Z}) \stackrel{d}{\approx} X, \quad \text{for a given } \mathbf{Z} \sim p_Z$$

- Ingredients:
 - The observations $\mathbf{x}_1, \dots, \mathbf{x}_n$ with their underlying properties.
 - Inputs: G_{θ} and the latent distribution p_Z .
 - The distance or the similarity criterion between p_{θ} and p_X , and the optimization process that defines the optimal θ .

Generative Modeling Framework

- Focus on a **parametric** family of generators $\{G_{\theta}\}_{\theta \in \Theta}$ where $\Theta \subset \mathbb{R}^p$ such that $G_{\theta}(\mathbf{Z}) \sim P_{\theta}$.
- The problem comes in finding the **'best' parameter** θ^* such that p_{θ} , and p_X are as close as possible:

$$G_{\theta^*}(\mathbf{Z}) \stackrel{d}{\approx} X, \quad \text{for a given } \mathbf{Z} \sim p_Z$$

- Ingredients:
 - The observations $\mathbf{x}_1, \dots, \mathbf{x}_n$ with their underlying properties.
 - Inputs: G_{θ} and the latent distribution p_Z .
 - The distance or the similarity criterion between p_{θ} and p_X , and the optimization process that defines the optimal θ .

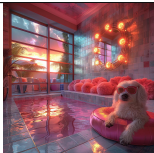
Generative Modeling Framework

- Focus on a **parametric** family of generators $\{G_{\theta}\}_{\theta \in \Theta}$ where $\Theta \subset \mathbb{R}^p$ such that $G_{\theta}(\mathbf{Z}) \sim P_{\theta}$.
- The problem comes in finding the **'best' parameter** θ^* such that p_{θ} , and p_X are as close as possible:

$$G_{\theta^*}(\mathbf{Z}) \stackrel{d}{\approx} X, \quad \text{for a given } \mathbf{Z} \sim p_Z$$

- Ingredients:
 - The observations $\mathbf{x}_1, \dots, \mathbf{x}_n$ with their underlying properties.
 - Inputs: G_{θ} and the latent distribution p_Z .
 - The distance or the similarity criterion between p_{θ} and p_X , and the optimization process that defines the optimal θ .

How AI Uses Generative Models

Model Type	Example Models	Applications	Image
Variational Autoencoders (VAEs)	VAE, β -VAE	Image reconstruction Anomaly detection Medical imaging	
Generative Adversarial Networks (GANs)	StyleGAN CycleGAN	Realistic face generation Image-to-image translation Deepfakes	
Diffusion Models	DALL·E 2 Imagen Midjourney	Text-to-image generation High-resolution artwork	

Neural networks

- **Neural networks:** Building block for generative models.
- Consider a neural network parametrization of G_θ .

Neural networks:

- Non-linear function built with a fixed number of neurons, each one representing a function, and distributed across several hidden layers.
- **One-hidden layer** neural network $G_\theta : \mathbb{R}^q \rightarrow \mathbb{R}$ composed of K **neurons**

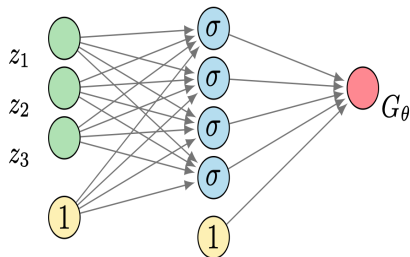
$$G_\theta(\mathbf{z}) = b_2 + \sum_{k=1}^K w_{2,k} \sigma(\langle \mathbf{w}_{1,k}, \mathbf{z} \rangle + b_{1,k}),$$

$$\theta := \{b_2\} \cup \{\mathbf{w}_{1,k}, w_{2,k}, b_{1,k}\}_{k=1}^K \in \Theta := \mathbb{R} \times (\mathbb{R}^q \times \mathbb{R} \times \mathbb{R})^K$$

where $\langle \cdot, \cdot \rangle$ is a scalar product on $\mathbb{R}^q$, and $\sigma : \mathbb{R} \rightarrow \mathbb{R}$ is a non-linear **activation function**, q is the dimension of $\mathbf{Z}$.

Neural networks

One-hidden layer neural network



$K = 4$ neurons and input dimension $q = 3$, the arrows are different parameters θ .

Figure from Handbook on Statistics of Extremes

Examples of activation functions:

- The cosine squasher

$$\sigma(x) = \frac{\cos(x + 3\pi/2) + 1}{2} I\left(x \in \left[-\frac{\pi}{2}, \frac{\pi}{2}\right]\right) + I\left(x \in \left(\frac{\pi}{2}, \infty\right)\right)$$

- The logistic squasher

$$\sigma(x) = \frac{1}{1 + e^{-x}}$$

- The exponential Linear Unit (eLU) defined for all $\alpha > 0$ by

$$\sigma_{\alpha}(x) = \alpha\{\exp(x) - 1\}I(x < 0) + xI(x \geq 0)$$

- The Rectified Linear Unit (ReLU)

$$\sigma(x) = \max(x, 0)$$

Neural networks

- Ideally $G_{\theta} \rightarrow G$ as $K \rightarrow \infty$ for different assumptions of G .

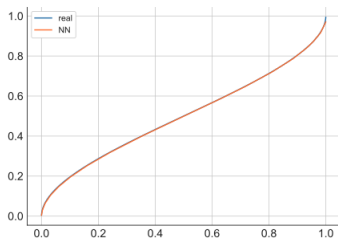
Theorem

(Universal approximation theorem). Suppose G is a continuous function on a compact space $\mathcal{Z} \subset \mathbb{R}^q$ and σ is not a polynomial. Then, for any $\varepsilon > 0$, there exists a one-hidden layer neural network G_{θ} (for some K depending on ε) such that

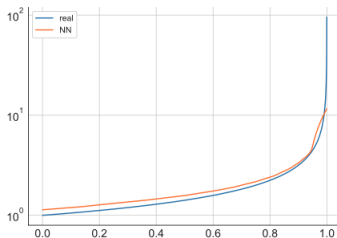
$$\sup_{\mathbf{z} \in \mathcal{Z}} |G(\mathbf{z}) - G_{\theta}(\mathbf{z})| < \varepsilon$$

- **Assumption:** function G to approximate is continuous on a compact set and thus **bounded**.

Examples $G = F_X^{-1}$ approximate the quantile function.



(1)



(2)

Figure: One-hidden layer neural network approximation of quantile functions $z \in [0, 1) \mapsto F_X^{-1}(z)$ from (a) Beta (2,2) and (b) Pareto ($\xi = 0.5$) distributions. The y-axis of (b) is in log-scale. *Figure from Handbook on Statistics of Extremes.*

Generative Adversarial Network

- **Goal:** Have a generator G such that $G(\mathbf{Z}) \overset{d}{\approx} X$, for $\mathbf{Z} \sim p_Z$
- A GAN aims at approximating the unknown generator G through a parametric **family of neural networks**

$$\{G_{\theta} : \mathbb{R}^q \rightarrow \mathbb{R}^D\}_{\theta \in \Theta}$$

- Estimate the optimal θ^* by optimizing an objective function: **adversarial game** between a **generator** and a **discriminator**.

$$\{D_{\phi} : \mathbb{R}^D \rightarrow [0, 1]\}_{\phi \in \Phi}$$

- Optimization problem:

$$\operatorname{argmin}_{\theta \in \Theta} \max_{\phi \in \Phi} \left(\mathbb{E}_{p_X} \{ \log D_{\phi}(\mathbf{X}) \} + \mathbb{E}_{p_Z} \{ \log [1 - D_{\phi} \{ G_{\theta}(\mathbf{z}) \}] \} \right)$$

Generative Adversarial Network

Learning about EDI from Data

- **Generator**: generate realistic data.
- **Discriminator**: tries to discriminate between synthetic generated data and real observations. $D_\phi(\mathbf{x})$ represents the probability that an observation $\mathbf{x} \in \mathbb{R}^D$ is drawn from $X \sim p_X$.

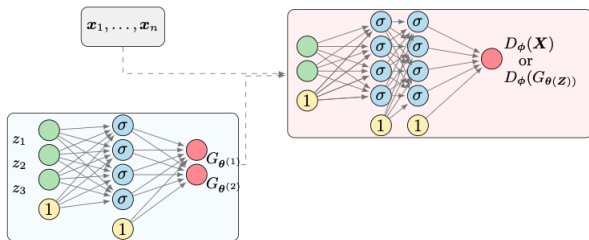


Figure: A GAN model with $q = 3$ (dimension of the latent noise) and $D = 2$ (dimension of data points). *Figure from Handbook on Statistics of Extremes.*

Main problem?

Simulating extremes

Theorem

(Universal approximation theorem). Suppose G is a **continuous function on a compact space** $\mathcal{Z} \subset \mathbb{R}^q$ and σ is not a polynomial. Then, for any $\varepsilon > 0$, there exists a one-hidden layer neural network G_θ (for some K depending on ε) such that

$$\sup_{\mathbf{z} \in \mathcal{Z}} |G(\mathbf{z}) - G_\theta(\mathbf{z})| < \varepsilon$$

- If G is **not bounded** there is no theoretical guarantee that a neural network could uniformly approximate G .
- **Generator**: $G(Z) \stackrel{d}{=} X$, for $Z \sim P_Z$
- Set $G := F_X^{-1}$ and $Z \sim U(0, 1)$.
- Since Z is a bounded random variable, when the activation function σ is continuous, $G_{\theta_K}(Z)$ is also a bounded random variable.

Why is this a problem?

Simulating extremes

- **Heavy tailed-distributions** (regularly varying $\mathcal{RV}_{-1/\gamma}$):

$$1 - F_X(x) = x^{-1/\gamma} \mathcal{L}_X(x), \quad x \rightarrow \infty$$

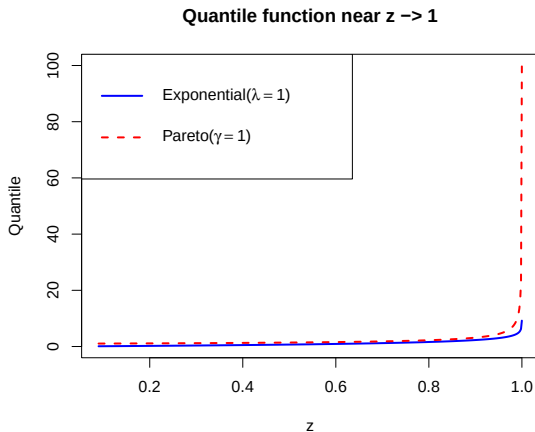
with $\mathcal{L}_X$ a slow varying function $\mathcal{L}(tx)/\mathcal{L}(x) \rightarrow 1$ for all $t > 0$.

- **Tail index** γ controls the tail: the larger the heavier.
- **Property of class $\mathcal{RV}$** : $F_X^{-1}(z) \sim (1 - z)^{-\gamma} \ell(1/(1 - z))$.
- And $F_X^{-1}(z) \rightarrow \infty$ as $z \rightarrow 1$ (fast).
- **Light-tailed**: Exponential distribution $1 - F(x) = e^{-\lambda x}$ and $F^{-1}(z) = -\log(1 - z)/\lambda$. $F_X^{-1}(z) \rightarrow \infty$ as $z \rightarrow 1$ (slow).

Why is this a problem?

Simulating extremes

- For light-tailed the quantile function diverges to infinity **slowly** compared to heavy tails.



Simulating extremes with GANs

Financial Illustration from Chapter

- Data: daily **log-returns** of financial stock market indices.
- Six **indices**: NKX (Nikkei, Japan), KOSPI (Korea), HSI (Hong-Kong), CAC (France), AMX (Amsterdam Exchange, Netherlands), Nasdaq (USA).
- GAN is assessed on:

Table: Tail index estimates from GAN approach applied to real financial data.

Index	Original Data	GAN Data
AEX	0.268	0.124
CAC	0.292	0.135
NKX	0.357	0.114
KOSPI	0.251	0.120
HSI	0.226	0.127
NDQ	0.352	0.166

Simulating extremes with GANs

Financial Illustration from Chapter

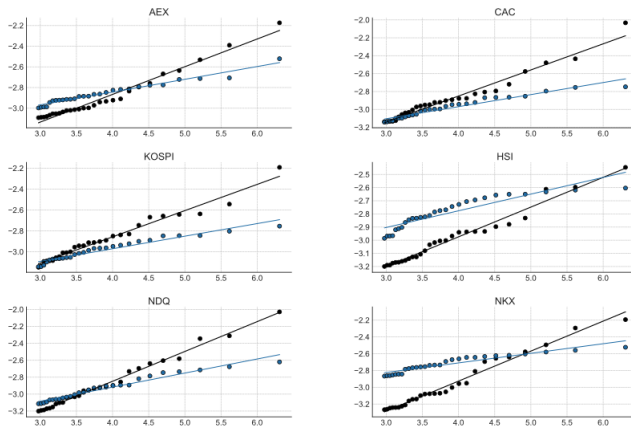


Figure: Log quantile-quantile plots from GAN approach $\xi = 0.95$ (black: real data, blue: GAN data). *Figure from Handbook on Statistics of Extremes.*

Improvements of GANs

- Some ideas:
 - Pre-processing of the data to get rid of the tail heaviness,
 - Use of heavy-tailed latent variables.
 - **Adapt the optimization problem to the heavy-tail situation.**
- EV-GAN and Tail-GAN

$$f^{\text{TIF}}(z) = -\frac{\log \{F_X^{-1}(z)\}}{\log \{\varphi(z)\}}, \quad \varphi(z) = \frac{1 - z^2}{2}, \quad z \in [0, 1)$$

is continuous, bounded on $[0, 1]$ and tends to the tail index as $z \rightarrow 1$.

- Theoretical **guarantee** that a neural network could uniformly approximate G .
- But usually not differentiable at $z = 1$.

EV-GAN

- Corrected version of the TIF is introduced:

$$f_{\boldsymbol{\beta}}^{\text{CTIF}}(z) = f^{\text{TIF}}(z) - \sum_{j=1}^6 \beta_j e_j(z)$$

where $\boldsymbol{\beta} = (\beta_1, \dots, \beta_6)^T$ is to be estimated, and $\{e_1, \dots, e_6\}$ are universal functions. When $\rho \in [-2, -1)$, there exist $\boldsymbol{\theta} \in \mathbb{R}^{3K+1}$, $\boldsymbol{\beta} \in \mathbb{R}^6$ and $C > 0$ such that

$$\sup_{z \in [0,1]} \left| f_{\boldsymbol{\beta}}^{\text{CTIF}}(z) - G_{\boldsymbol{\theta}}(z) \right| \leq C \times K^{\tau}$$

- Leading to

$$W_1(\boldsymbol{\theta}, \boldsymbol{\beta}) = \int_0^1 \left| F_X^{-1}(z) - \varphi(z)^{-\tilde{G}_{\boldsymbol{\theta}, \boldsymbol{\beta}}(z)} \right| dz$$

where $\tilde{G}_{\boldsymbol{\theta}, \boldsymbol{\beta}}$ is the enriched version of the generator $G_{\boldsymbol{\theta}}$ defined as

$$\tilde{G}_{\boldsymbol{\theta}, \boldsymbol{\beta}}(z) = G_{\boldsymbol{\theta}}(z) + \sum_{j=1}^6 \beta_j e_j(z)$$

EV-GAN

- The **optimization** process is the **same** as before.
- The only difference comes from the **generator's design**.
- EV-GAN Generator
 - 1 Require: Trained parameters (θ, β)
 - 2 Generate latent variable $\mathbf{Z}$ in dimension $q \geq D$, from $U(0, 1)$
 - 3 for each marginal $d = 1 : D$ do
 - 4 $G_{\theta, \beta}^{\text{EVGAN}, (d)}(\mathbf{Z}) = \varphi(Z_d)^{-\{G_{\theta(Z)} + \sum_{j=1}^6 \beta_j e_j Z_d\}}$
 - 5 end for
- The generator is:

$$X \stackrel{d}{\approx} \varphi(Z)^{-\tilde{G}_{\theta, \beta}(Z)}, \quad \text{with} \quad Z \stackrel{d}{=} U(0, 1).$$

Simulating extremes with GANs

Financial Illustration from Chapter

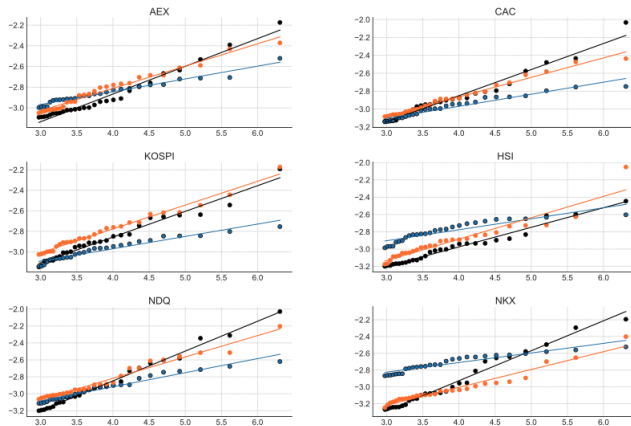


Figure: Log quantile-quantile plots from GAN approach $\xi = 0.95$ (black: real data, blue: GAN data, orange: EV-GAN). *Figure from Handbook on Statistics of Extremes.*

Closing Remarks

Final Comments

- When the **target random quantity** is supposed to be **heavy-tailed**, the usual generative models **cannot reproduce** this property without dedicated architecture improvements.
- For example EV-GAN.
- The chapter also discussed improvements for **Variational Auto-Encoders** and **Difussion Models**.

Based on Chapter from Handbook on Statistics of Extremes

This talk is based on:

Chapter 20 – “*On the Simulation of Extreme Events with Neural Networks*”

by Michaël Allouche, Stéphane Girard & Emmanuel Gobet

In *Handbook on Statistics of Extremes*

