



Extreme risk decision modeling: a new approach with an application to air pollution based on PM_{2.5}

Irene Brito ^{*}, Marta Ferreira 

Centro de Matemática, Departamento de Matemática, Universidade do Minho, 4800-058, Guimarães, Portugal

ARTICLE INFO

Keywords:

Risk decision model
Risk assessment
Risk measure
Extreme value theory
Generalized Pareto distribution
Air pollution

ABSTRACT

Extreme risks are events or scenarios that have a very low occurrence probability, however they can imply catastrophic consequences. These risks are often beyond the range of normal expectations and can lead, depending on the context, to financial or economic losses, loss of life, or environmental damage. The main purpose is to develop extreme risk decision models that account for the presence of extreme values. Two-component risk models based on mean and standard deviation and on mean and value at risk are defined for the Generalized Pareto distribution and properties are investigated in function of the model parameters. It is also shown that the new risk models can be represented as a multi-criteria risk assessment method: the additive ratio assessment method. The risk models are applied to air pollution risk assessment by analysing fine particulate matter (PM_{2.5}) concentrations in a European city, in order to identify the riskiest months and periods of the year.

1. Introduction

Risk measures and risk models have predominantly been developed and applied in economic, financial, or actuarial contexts to support decision making under risk. Some of them are built upon the expected utility theory formulated by von Neumann and Morgenstern (1944). According to this theory, when confronted with uncertain outcomes, individuals seek to maximize their expected utility rather than their expected monetary value. Because expected utility theory often fails to align with empirical findings – as noted in Kahneman and Tversky (1979) – extensions and alternatives to the model have been proposed. Some generalizations combine expected utility with other risk measures (see e.g. Geissel et al. (2022), Brito (2020), Brito (2022)), including for example variance or entropy. These belong to the naive risk measures (see Brachinger and Weber (1997)), that can be applied in pure or combined form to risk decision problems (see also Jia and Dyer (2009)). The mean and variance are the most commonly used measures (see for instance Markowitz (2000), Pollatsek and Tversky (1970)). Examples of extended risk models that incorporate additional measures are the mean-variance-skewness model (Li et al., 2010), the mean-variance-skewness-entropy model (Usta and Kantar, 2011), the mean-variance-skewness-kurtosis-entropy model (Aksarayli and Pala, 2018). Other popular risk measures that are widely used in the financial or actuarial context are the loss probability or the value at risk (Kaas et al., 2008; Klugman et al., 2019), that represents a maximum value (an adverse value) that will not be exceeded with a certain probability.

In environmental risk assessment, in particular in the context of air pollution, the methods and techniques developed and proposed in the literature consist mainly of probabilistic approaches, focusing on the analysis of mean pollutant concentrations, exceedances

^{*} Corresponding author.

E-mail addresses: ireneb@math.uminho.pt (I. Brito), msferreira@math.uminho.pt (M. Ferreira).

of regulatory limits, and air quality indices. The main shortcoming of those methods is their inability to account for specific risk characteristics, such as uncertainty and the presence of extreme values.

The aim of the present work is to develop risk models that assess risk using two components: mean combined with standard deviation, the Mean-SD Risk Model, and mean combined with value at risk, the Mean-VaR Risk Model, integrated through a trade-off factor. Furthermore, these models explicitly account for the distinct characteristics of extreme values, since they are constructed using the Generalized Pareto (GP) distribution within the framework of extreme value theory. The GP distribution is primarily used in fields that deal with rare and extreme events, where understanding the tail of a distribution is crucial (Balkema and Haan, 1974; Pickands, 1975; Davison and Smith, 1990; Beirlant et al., 2004). From environmental studies and risk management to finance and insurance, the GP provides valuable insights into the likelihood of extreme outcomes, helping organizations plan and mitigate risks effectively (Castillo and Serra, 2015; Lomba and Alves, 2020; Moreira and Ferreira, 2025). To rigorously characterize the probability and severity of extreme values, i.e. those observations exceeding a high threshold, one typically turns to Extreme Value Theory. A cornerstone result in this area is the Pickands-Balkema-de Haan theorem (Balkema and Haan, 1974; Pickands, 1975), which asserts that for a sufficiently large threshold (μ), the conditional distribution of exceedances approaches a GP distribution. The GP is parametrized by a scale parameter ($\sigma > 0$) and a shape parameter ($\gamma \in \mathbb{R}$) also denoted extreme value index or tail index. This flexibility makes it ideal for the “peaks over threshold” approach: one selects a high threshold (μ), fits the GP to the observed exceedances (e.g. by maximum likelihood), and then uses the fitted model to extrapolate the probabilities and magnitudes of even more extreme events.

Regarding the two-component risk models, it will be shown that these can be presented as a multi-criteria decision making (MCDM) method Sahoo and Goswami (2023): additive ratio assessment (ARAS) method (Zavadskas and Turskis, 2010) (see also Liu and Xu (2021)). In the ARAS method, alternatives are evaluated by first normalizing and weighting each criterion and then aggregating them to compute a utility score that indicates how close each alternative is to the ideal solution. The alternatives can then be ranked according to their utility scores. In order to reformulate the risk models as ARAS methods, the normalization technique of the ARAS method will be adopted to define the normalized risk models proposed here.

The new risk models are applicable to a variety of contexts involving extreme events. This study focuses on their application to air pollution risk assessment based on fine particulate matter ($PM_{2.5}$) concentrations.

Air pollution is one of the largest environmental health risk. Numerous studies have shown that exposure to air pollution negatively impacts health, leading to respiratory problems, cardiovascular diseases, and reduced life expectancy (see e.g. Orellano et al. (2020)). The air pollutant fine particulate matter ($PM_{2.5}$) is considered the most harmful air pollutant globally. $PM_{2.5}$ stands for particulate matter with a diameter of $2.5\mu m$ or less. These particles are emitted mainly from the combustion of solid fuels for domestic heating, industrial activities and road transport. In the European Union (EU), despite ongoing reductions in emissions and overall improvements in air quality, most of the EU’s urban population continued to be exposed to levels of key air pollutants that are damaging to health, and 96% of the urban population is exposed to unsafe concentrations of $PM_{2.5}$ (see EEA (2024)), above the 2021 World Health Organization (WHO) annual guideline level of $5\mu g/m^3$. Regarding the EU standards related to the Average Exposure Indicator for $PM_{2.5}$ (see EEA (2024)), which assesses the general population’s long-term exposure in urban areas, all EU Member States continued to meet the exposure concentration obligation of $20\mu g/m^3$ in 2022, set as a 2015 target under the ambient air quality directive.

Examples of articles dealing with quantitative air pollution risk assessment of $PM_{2.5}$ that appeared in the literature are the following. In Liu et al. (2023), the authors analysed the $PM_{2.5}$ winter levels in Beijing in order to predict and estimate the risk for air pollution based on average concentrations. In Lin et al. (2022), an air quality framework was proposed to determine air quality zones in Taiwan using $PM_{2.5}$ risk sensitivity for exceedance probability. Extreme value analysis of atmospheric pollutant data, including the application of the GP distribution to elevated concentrations of $PM_{2.5}$ (among other air quality indicators), was presented in Martins et al. (2017). The application of extreme value models to pollution data - specifically, to the modeling of high ozone concentration levels - was also considered in Ferreira et al. (2012) and Leiva et al. (2016).

In the context of $PM_{2.5}$ air pollution analysis, the novelty of this work is that it uses a quantitative risk assessment approach that incorporates combined risk measures, while also integrating extreme value theory into the risk modeling framework.

The paper is organized as follows. In Section 2, the risk measures and associated risk models are defined and their properties are investigated, where Section 2.1 is dedicated to the Mean-standard deviation Risk Model and to the Mean-Value at Risk Model and Section 2.2 presents the properties of those models. Section 3 introduces the extreme value theory approach based on the Generalized Pareto distribution to establish the extreme risk measures and the corresponding extreme risk models: the Generalized Pareto Mean-standard deviation Risk Model and the Generalized Pareto Mean-Value at Risk Model. The properties of these risk models are analysed based on their parameters. Section 4 demonstrates that the new risk models can be represented as a multi-criteria decision making method: the additive ratio assessment (ARAS) method. Section 5 presents the application of the extreme risk models to air pollution risk assessment caused by $PM_{2.5}$ in a European city, where the aim is to classify months and periods of the year in terms of the risk for air pollution. The conclusions are presented in Section 6.

2. Risk models

Pollatsek and Tversky (1970) formulated the following risk model. Given a set S of real-valued random variables, interpreted as gambles with (arbitrary numbers of) monetary outcomes, for $X \in S$ with expectation $E[X]$ and variance $\text{var}[X]$ its risk $R(X)$ is given by

$$R(X) = \theta \text{var}[X] - (1 - \theta)E[X], \quad (1)$$

where θ is a real number, such that $0 < \theta \leq 1$. For any $X, Y \in S$, $Y \leq X$ means that X is at least as risky as Y and the risk ordering can be determined by $Y \leq X \Leftrightarrow R(Y) \leq R(X)$. According to this risk model, the risk ordering is generated by a linear combination of expectation and variance. The parameter θ specifies the relative contribution of the expectation and the variance to the riskiness of the lottery.

Usually, for risk assessment purposes in different contexts, e.g. in the environmental context, risk indices or risk factors are calculated, mainly as the product of an occurrence probability and a consequence, that represent expected values. However, as [Aven \(2016\)](#) pointed out, after analysing several works, risk assessment that relies on principles such as the general conception of risk as an expected value, could misguide decision-makers, and there is the need for taking also into account uncertainties.

Motivated by the mean-variance model we will develop two risk models, based on the following assumption, with all subsequent results being presented in this context.

Assumption 1. The random variable X is a risk related random variable, where higher outcomes of X represent higher risks.

2.1. Mean-SD risk model and Mean-VaR risk model

In the context of air pollution risk assessment, where the random variable X represents the concentration (or measurement) of a certain air pollutant and assuming that risk for environmental and human health increases with higher concentrations of that pollutant and, thus, [Assumption 1](#) holds, risk increases with expectation, however also with variance, due to the higher uncertainty related with higher dispersion. According to this risk perception, the following risk measure is defined, maintaining the structure of the linear combination as in the mean-variance model, where, however, for dimension coherence, the standard deviation is preferred over the variance.

Definition 1 (Mean-SD risk measure). The Mean-SD measure of risk for X is defined by

$$R_1(X, \theta) = \theta \text{SD}[X] + (1 - \theta)E[X], \quad (2)$$

where θ is a real constant satisfying $0 \leq \theta \leq 1$.

The trade-off parameter θ combines the contribution of both risk components, expectation and standard deviation, to the final risk and can be used to analyse the impact of each risk component on the final risk. If $\theta = 0$, then risk is judged only by the expected value and if $\theta = 1$, then risk is assessed only by standard deviation. For $\theta = \frac{1}{2}$ the risk components contribute equally to the final risk. If $\theta \in (0, 1)$, then the effect of the expectation on the risk measure is higher if θ approaches 0, and if θ tends towards 1, then the uncertainty has a greater influence on the risk measure.

In order to set up a risk decision model based on this risk measure for the assessment of different alternatives whose associated risks can be presented by different random variables, we define the following normalized Mean-SD risk measure. The normalization of the risk factors helps improving the relative contribution of each factor through the aggregation and ensures that the overall risks are interpreted in a consistent way and are comparable for the different random variables.

Definition 2 (Normalized Mean-SD risk measure). Consider the set of random variables $\mathbb{X} = \{X_1, \dots, X_I\}$ and let $X \in \mathbb{X}$. The normalized Mean-SD measure of risk for X is defined by

$$\bar{R}_1(X, \theta) = \theta \bar{\text{SD}}[X] + (1 - \theta)\bar{E}[X], \quad (3)$$

where θ is a real constant satisfying $0 \leq \theta \leq 1$ and $\bar{\text{SD}}[X]$ and $\bar{E}[X]$ are the normalized standard deviation and normalized mean, respectively, defined by

$$\bar{\text{SD}}[X] = \frac{\text{SD}[X]}{\sum_{i=1}^{I+1} \text{SD}[X_i]}, \quad \text{where } \text{SD}[X_{I+1}] = \min_{X_i \in \mathbb{X}} \{\text{SD}[X_i]\},$$

$$\bar{E}[X] = \frac{E[X]}{\sum_{i=1}^{I+1} E[X_i]}, \quad \text{where } E[X_{I+1}] = \min_{X_i \in \mathbb{X}} \{E[X_i]\},$$

and $\text{SD}[X]$ and $E[X]$ are, respectively the standard deviation and mean of X .

Regarding the normalization, here the normalization by division through the sum - including the minimum value - is adopted, since this is the normalization used in the ARAS method ([Zavadskas and Turskis, 2010](#)). As it will be seen later, the risk model can in this way be related with the ARAS method, where the additional minimum value corresponds to the optimal value. The optimal value represents in this case the minimum risk (cf. [Brito \(2025\)](#)).

The Value at Risk is a widely adopted risk measure, especially in the financial and actuarial context, for understanding and managing financial risk, where it is used to quantify potential losses. Given a random variable X , the VaR at confidence level $1 - p$, denoted by $\text{VaR}[X; p]$, is defined as

$$\text{VaR}[X; p] = \inf\{x \in \mathbb{R} : P(X \leq x) \geq 1 - p\}, \quad (4)$$

where $p = 0.05$ and $p = 0.01$ will be used. The VaR risk measure is also appropriate for applications in other contexts, such as in environmental risk assessment. In air pollution risk assessment, where the random variable X represents the concentration of an air pollutant and assuming that the risk of air pollution increases when the concentration increases, $\text{VaR}[X; p]$ defines the worst pollutant concentration at the confidence level $1 - p$ and risk increases with $\text{VaR}[X; p]$. Moreover, $\text{VaR}[X; p]$ can be interpreted as a return level, a concept commonly used in extreme-value theory, corresponding in this case to the pollutant concentration expected to be exceeded on average once every $1/p$ observations (or time periods). This motivates the definition of the following risk measure.

Definition 3 (Mean-VaR risk measure). The Mean-VaR measure of risk for X is defined by

$$R_2(X, \theta) = \theta \text{VaR}[X; p] + (1 - \theta) \mathbb{E}[X], \quad (5)$$

where θ is a real constant satisfying $0 \leq \theta \leq 1$.

For the risk assessment involving different random variables the following normalized risk measure will be used.

Definition 4 (Normalized Mean-VaR risk measure). Consider the set of random variables $\mathbb{X} = \{X_1, \dots, X_I\}$ and let $X \in \mathbb{X}$. The normalized Mean-VaR measure of risk for X is defined by

$$\bar{R}_2(X, \theta) = \theta \bar{\text{VaR}}[X; p] + (1 - \theta) \bar{\mathbb{E}}[X], \quad (6)$$

where θ is a real constant satisfying $0 \leq \theta \leq 1$ and $\bar{\text{VaR}}[X; p]$ and $\bar{\mathbb{E}}[X]$ are the normalized value at risk and normalized mean, respectively, defined by

$$\bar{\text{VaR}}[X; p] = \frac{\text{VaR}[X; p]}{\sum_{i=1}^{I+1} \text{VaR}[X_i; p]}, \quad \text{where } \text{VaR}[X_{I+1}; p] = \min_{X_i \in \mathbb{X}} \{\text{VaR}[X_i; p]\},$$

$$\bar{\mathbb{E}}[X] = \frac{\mathbb{E}[X]}{\sum_{i=1}^{I+1} \mathbb{E}[X_i]}, \quad \text{where } \mathbb{E}[X_{I+1}] = \min_{X_i \in \mathbb{X}} \{\mathbb{E}[X_i]\},$$

and $\text{VaR}[X; p]$ and $\mathbb{E}[X]$ are, respectively the value at risk and mean of X .

Based on these risk measures, we define the risk order relation in [Definition 5](#) for assessing the risk for different alternatives. In the present context, the alternatives, denoted by A_i , represent different elements or options that are evaluated, each associated with a risk random variable, X_i . For example, in environmental risk assessment, the alternatives may represent different locations and the associated random variables describe the concentration of a specific pollutant at each location. Or, the alternatives may represent different time intervals, and the random variables describe the pollutant concentration in the different time intervals at a single location.

Definition 5. Let $A_1, \dots, A_I$ be different alternatives with corresponding random variables $X_1, \dots, X_I$.

- (a) Consider two alternatives A_1 and A_2 with corresponding normalized Mean-SD risk measures $\bar{R}_1(X_1, \theta)$ and $\bar{R}_1(X_2, \theta)$, defined in (3), or normalized Mean-VaR risk measures $\bar{R}_2(X_1, \theta)$ and $\bar{R}_2(X_2, \theta)$, defined in (6). Then:
 - (i) A_1 is riskier than A_2 , $A_1 > A_2$, if $\bar{R}_i(X_1, \theta) > \bar{R}_i(X_2, \theta)$, $i = 1, 2$.
 - (ii) A_1 is weakly riskier than A_2 , $A_1 \succeq A_2$, if $\bar{R}_i(X_1, \theta) \geq \bar{R}_i(X_2, \theta)$, $i = 1, 2$.
 - (iii) A_1 is equally as risky as A_2 , $A_1 \sim A_2$, if $\bar{R}_i(X_1, \theta) = \bar{R}_i(X_2, \theta)$, $i = 1, 2$.
- (b) Consider multiple alternatives, $A_1, A_2, \dots, A_I$. Then, the normalized risk measures permit ranking the alternatives, where the optimal alternative is the one that minimizes the risk measure, so that A_l , $l \in \{1, \dots, I\}$, is the least risky alternative if

$$\bar{R}_i(X_l, \theta) = \min_{k \in \{1, \dots, I\}} \bar{R}_i(X_k, \theta),$$

$i = 1, 2$.

2.2. Properties of the risk models

The following results show that the proposed risk models are adequate for the risk description of environmental air pollution assessment problems.

The risk decisions are straightforward in the following cases.

Proposition 1. Let X_1 and X_2 be risk random variables.

1. Consider the normalized Mean-SD risk measure defined in (3). Then:
 - (a) If $\mathbb{E}[X_1] > \mathbb{E}[X_2]$ and $\text{SD}[X_1] > \text{SD}[X_2]$, then $\bar{R}_1(X_1, \theta) > \bar{R}_1(X_2, \theta)$.
 - (b) If $\mathbb{E}[X_1] < \mathbb{E}[X_2]$ and $\text{SD}[X_1] < \text{SD}[X_2]$, then $\bar{R}_1(X_1, \theta) < \bar{R}_1(X_2, \theta)$.
2. Consider the normalized Mean-VaR risk measure defined in (6). Then:
 - (a) If $\mathbb{E}[X_1] > \mathbb{E}[X_2]$ and $\text{VaR}[X_1; p] > \text{VaR}[X_2; p]$, then $\bar{R}_2(X_1, \theta) > \bar{R}_2(X_2, \theta)$.
 - (b) If $\mathbb{E}[X_1] < \mathbb{E}[X_2]$ and $\text{VaR}[X_1; p] < \text{VaR}[X_2; p]$, then $\bar{R}_2(X_1, \theta) < \bar{R}_2(X_2, \theta)$.

The properties in [Proposition 2](#) follow directly from the definition and indicate that risk increases (decreases) with higher (lower) values of the random variable, either through the addition of a positive (negative) constant or through the multiplication of a constant greater than 1 (constant belonging to the interval $(0, 1)$). These properties are in agreement with the perception of people about the description of risk in the context of air pollution, where increasing (decreasing) additive or multiplicative air pollutant concentrations imply a higher (lower) risk.

Proposition 2. Let X be a risk random variable and consider the normalized Mean-SD risk measure $\bar{R}_1$, defined in (3), and the normalized Mean-VaR risk measure $\bar{R}_2$, defined in (6).

1. Consider $\mathbb{X} = \{X, X + \alpha\}$, with $\alpha \in \mathbb{R}$. Then:

$$\bar{R}_i(X + \alpha, \theta) > \bar{R}_i(X, \theta), \text{ if } \alpha > 0, \quad (7)$$

$$\bar{R}_i(X + \alpha, \theta) < \bar{R}_i(X, \theta), \text{ if } \alpha < 0, \quad (8)$$

$i = 1, 2$.

2. Consider $\mathbb{X} = \{X, \alpha X\}$, with $\alpha \in \mathbb{R}^+$. Then:

$$\bar{R}_i(\alpha X, \theta) > \bar{R}_i(X, \theta), \text{ if } \alpha > 1, \quad (9)$$

$$\bar{R}_i(\alpha X, \theta) < \bar{R}_i(X, \theta), \text{ if } 0 < \alpha < 1, \quad (10)$$

$i = 1, 2$.

Proof. The proofs of this and the following propositions are provided in [Appendix A](#). $\square$

Proposition 3 demonstrates that the risk order relation is invariant under changes in scale. Furthermore, if one scale results in a lower risk than another for a specific random variable, this risk ordering will persist for all other random variables evaluated using those scales.

Proposition 3. Consider the normalized Mean-SD risk measure $\bar{R}_1$, defined in (3), the normalized Mean-VaR risk measure $\bar{R}_2$, defined in (6), and two alternatives with risk random variables X_1 and X_2 , and $\alpha, \beta \in \mathbb{R}^+$ with $\alpha \neq \beta$.

1. Let $\mathbb{X} = \{X_1, X_2\}$. Then, $\bar{R}_i(X_1, \theta) < \bar{R}_i(X_2, \theta)$ if and only if $\bar{R}_i(\alpha X_1, \theta) < \bar{R}_i(\alpha X_2, \theta)$, $i = 1, 2$, for $\mathbb{X} = \{\alpha X_1, \alpha X_2\}$.

2. Let $\mathbb{X} = \{\alpha X_1, \beta X_1\}$. Then $\bar{R}_i(\alpha X_1, \theta) < \bar{R}_i(\beta X_1, \theta)$ if and only if $\bar{R}_i(\alpha X_2, \theta) < \bar{R}_i(\beta X_2, \theta)$, $i = 1, 2$, for $\mathbb{X} = \{\alpha X_2, \beta X_2\}$.

3. Generalized Pareto risk models

In extreme value theory, the Generalized Pareto (GP) distribution is used to model the tails of distributions, especially when dealing with extreme events. The GP model is able to describe the behaviour of the largest values in a given dataset or in the tail, e.g. such as the highest pollutant concentrations. More precisely, fixing a sufficiently high threshold μ for a given random variable Y , the conditional excesses $X|X > 0$, where $X = Y - \mu$, can be well approximated by a GP model defined by the following cumulative distribution function:

$$H(x) = \begin{cases} 1 - \left(1 + \gamma \frac{x}{\sigma}\right)_+^{-\frac{1}{\gamma}} & , \quad \gamma \neq 0 \\ 1 - \exp\left(-\frac{x}{\sigma}\right) & , \quad \gamma = 0, \end{cases} \quad (11)$$

with $(z)_+ = \max(0, z)$, where $\sigma > 0$ is the scale parameter and γ is the shape parameter ([Pickands, 1975](#); [Balkema and Haan, 1974](#)).

If $\gamma > 0$, the distribution has a heavy tail, if $\gamma = 0$, the distribution simplifies to an exponential tail, and, if $\gamma < 0$, the distribution has a bounded tail.

The mean of the GP distribution exists when $\gamma < 1$ and is given by

$$\mathbb{E}[X] = \frac{\sigma}{1 - \gamma}. \quad (12)$$

The variance of the GP distribution exists when $\gamma < \frac{1}{2}$ and is expressed by

$$\text{var}[X] = \frac{\sigma^2}{(1 - \gamma)^2(1 - 2\gamma)}. \quad (13)$$

The value at risk $\text{VaR}[X; p]$ of the GP distribution is determined by

$$\text{VaR}[X; p] = \begin{cases} \frac{\sigma}{\gamma}(p^{-\gamma} - 1) & , \quad \gamma \neq 0 \\ -\sigma \log p & , \quad \gamma = 0. \end{cases} \quad (14)$$

Replacing (12) and (13) in the expression (2) of [Definition 1](#), we define the Generalized Pareto Mean-SD Risk Model, where the risk decisions are taken according to the risk order relation described in [Definition 5](#).

Definition 6 (GP Mean-SD risk measure). Let X be a random variable following a GP distribution defined in (11) with parameters $\sigma > 0$ and $\gamma < \frac{1}{2}$. The GP Mean-SD measure of risk for X is defined by

$$R_1^{\text{GP}}(\gamma, \sigma, \theta) = \theta \frac{\sigma}{(1 - \gamma)\sqrt{1 - 2\gamma}} + (1 - \theta) \frac{\sigma}{1 - \gamma}, \quad (15)$$

where θ is a real constant satisfying $0 \leq \theta \leq 1$.

We will now analyse the GP-Mean-SD measure of risk in function of its parameters γ , σ and θ .

Proposition 4. The GP-Mean-SD measure of risk $R_1^{GP}(\gamma, \sigma, \theta)$ defined in (15) increases with σ ; increases with γ ; increases with θ if $0 < \gamma < \frac{1}{2}$, decreases with θ if $\gamma < 0$, and $R_1^{GP}(\gamma, \sigma, \theta)$ is constant in θ if $\gamma = 0$, in which case $R_1^{GP}(\gamma, \sigma, \theta) = \sigma$.

The results in the previous proposition align with risk perception: a higher scale parameter indicates a greater mean and standard deviation, while a higher shape parameter implies a heavier tail, and consequently, the risk increases with these parameters. Regarding the role of the trade-off parameter, note that the risk measure (15) can be expressed as

$$R_1^{GP}(\gamma, \theta) = \theta \frac{\mathbb{E}[X]}{\sqrt{1-2\gamma}} + (1-\theta)\mathbb{E}[X] = \theta \left(\frac{\mathbb{E}[X]}{\sqrt{1-2\gamma}} - \mathbb{E}[X] \right) + \mathbb{E}[X].$$

For $\theta = 0$ the risk equals $\mathbb{E}[X]$. Increasing θ increases the risk if $0 < \gamma < \frac{1}{2}$ by adding the term $\theta \frac{\mathbb{E}[X]}{\sqrt{1-2\gamma}} - \mathbb{E}[X]$ to $\mathbb{E}[X]$, since in this case the term $\frac{\mathbb{E}[X]}{\sqrt{1-2\gamma}} - \mathbb{E}[X]$ is positive, meaning that the uncertainty component is higher than the mean, $\text{SD}[X] > \mathbb{E}[X]$. However, if $\gamma < 0$, increasing θ decreases the risk, since in this case the term $\frac{\mathbb{E}[X]}{\sqrt{1-2\gamma}} - \mathbb{E}[X]$ is negative and consequently the term $\theta \frac{\mathbb{E}[X]}{\sqrt{1-2\gamma}} - \mathbb{E}[X]$ will be subtracted from $\mathbb{E}[X]$, in which case $\text{SD}[X] < \mathbb{E}[X]$.

For the risk assessment and comparison of different random variables, the following normalized GP risk measure is defined.

Definition 7 (Normalized GP Mean-SD risk measure). Consider a set of risk random variables $\mathbb{X} = \{X_1, \dots, X_I\}$. Let $X_k \in \mathbb{X}$ be a random variable with GP distribution defined in (11), with parameters $\sigma_k > 0$ and $\gamma_k < \frac{1}{2}$, and let $X_i, i = 1, \dots, I$, be random variables with GP distributions, with parameters $\sigma_i > 0$ and $\gamma_i < \frac{1}{2}$. The normalized GP Mean-SD measure of risk for X_k is defined by

$$\bar{R}_1^{GP}(\gamma_k, \sigma_k, \theta) = \frac{\theta}{1 + \frac{(1-\gamma_k)\sqrt{1-2\gamma_k}}{\sigma_k} \sum_{\substack{i=1 \\ i \neq k}}^{I+1} \frac{\sigma_i}{(1-\gamma_i)\sqrt{1-2\gamma_i}}} + \frac{(1-\theta)}{1 + \frac{1-\gamma_k}{\sigma_k} \sum_{\substack{i=1 \\ i \neq k}}^{I+1} \frac{\sigma_i}{1-\gamma_i}}, \quad (16)$$

where θ is a real constant satisfying $0 \leq \theta \leq 1$ and

$$\frac{\sigma_{I+1}}{(1-\gamma_{I+1})\sqrt{1-2\gamma_{I+1}}} = \min_{i=1, \dots, I} \left\{ \frac{\sigma_i}{(1-\gamma_i)\sqrt{1-2\gamma_i}} \right\},$$

$$\frac{\sigma_{I+1}}{1-\gamma_{I+1}} = \min_{i=1, \dots, I} \left\{ \frac{\sigma_i}{1-\gamma_i} \right\}.$$

Proposition 5. Consider the normalized GP-Mean-SD measure of risk $\bar{R}_1^{GP}(\gamma_k, \sigma_k, \theta)$ defined in (16) and let

$$F(\gamma_k) = \sum_{\substack{i=1 \\ i \neq k}}^{I+1} \frac{\sigma_i(\sqrt{1-2\gamma_i} - \sqrt{1-2\gamma_k})}{(1-\gamma_i)\sqrt{1-2\gamma_i}}.$$

Then, $\bar{R}_1^{GP}(\gamma_k, \sigma_k, \theta)$ increases with σ_k ; increases with γ_k ; increases with θ if $F(\gamma_k) > 0$, decreases with θ if $F(\gamma_k) < 0$, and $\bar{R}_1^{GP}(\gamma_k, \sigma_k, \theta)$ is constant in θ if $F(\gamma_k) = 0$.

As a consequence of the previous proposition, in the following situations, for an alternative with maximum or minimum tail index, one can conclude whether the risk measure increases or decreases with θ .

Corollary 1. Let $\mathbb{X} = \{X_1, \dots, X_I\}$, $X_k \in \mathbb{X}$ and consider the normalized GP-Mean-SD measure of risk $\bar{R}_1^{GP}(\gamma_k, \sigma_k, \theta)$ in Definition(7). If $\gamma_k > \gamma_i$ for all $i = 1, \dots, I$ with $i \neq k$, then $\bar{R}_1^{GP}(\gamma_k, \sigma_k, \theta)$ increases with θ , if $\gamma_k < \gamma_i$ for all $i = 1, \dots, I$ with $i \neq k$, then $\bar{R}_1^{GP}(\gamma_k, \sigma_k, \theta)$ decreases with θ , and if $\gamma_k = \gamma_i, i = 1, \dots, I$, then $\bar{R}_1^{GP}(\gamma_k, \sigma_k, \theta)$ remains constant in θ .

Thus, given a set of alternatives, risk increases with θ for the alternative with maximum tail index, since $F(\gamma_k) > 0$, whereas risk decreases with θ for the alternative with minimum tail index, due to the fact that $F(\gamma_k) < 0$. This means that increasing the trade-off factor penalizes alternatives with maximum tail and favors alternatives with minimum tail by lowering the risk. If the alternatives have identical tail indices, the trade-off factor has no effect on the risk. In that case, only the scale parameter influences the risk of the alternatives.

The Generalized Pareto Mean-VaR risk model is obtained by substituting (12) and (14) into (5) of Definition 3 and the risk decisions are made according to the risk order relation described in Definition 5. Note that in the following results the condition that the shape parameter be less than 0.5 is maintained, since this condition ensures that all three risk components (mean, standard deviation and value at risk), and thus both risk models (Mean-SD and Mean-VaR risk models) are well defined.

Definition 8 (GP Mean-VaR risk measure). Let X be a random variable following a GP distribution defined in (11). The GP Mean-VaR measure of risk for X is defined by

$$R_2^{GP}(\gamma, \sigma, \theta; p) = \begin{cases} \frac{\theta\sigma}{\gamma}(p^{-\gamma} - 1) + \frac{(1-\theta)\sigma}{1-\gamma}, & \gamma \neq 0 \\ \theta(-\sigma \log p) + (1-\theta)\sigma, & \gamma = 0, \end{cases} \quad (17)$$

where θ is a real constant satisfying $0 \leq \theta \leq 1$.

Analysing this risk measure as a function of its parameters yields the following result.

Proposition 6. The GP Mean-VaR measure of risk $R_2^{GP}(\gamma, \sigma, \theta; p)$ defined in (17) increases with σ , increases with γ and increases with θ for $p \in (0, 0.05]$.

The results in Proposition 6 demonstrate that risk increases with the shape parameter, leading to a higher mean and a higher VaR. Additionally, risk rises with a heavier tail and when more weight is assigned to the VaR component.

Definition 9 (Normalized GP Mean-VaR risk measure).

Consider a set of risk random variables $\mathbb{X} = \{X_1, \dots, X_I\}$. Let $X_k \in \mathbb{X}$ be a random variable following a GP distribution defined in (11) with parameters σ_k and γ_k , and let $X_i, i = 1, \dots, I$, be random variables following GP distributions with parameters σ_i and γ_i . The normalized GP Mean-VaR measure of risk for X_k is defined by

$$\bar{R}_2^{GP}(\gamma_k, \sigma_k, \theta; p) = \begin{cases} \frac{\theta}{1 + \frac{\gamma_k}{\sigma_k(p^{-\gamma_k} - 1)}\Pi} + \frac{(1 - \theta)}{1 + \frac{1 - \gamma_k}{\sigma_k} \sum_{i=1, i \neq k}^{I+1} \frac{\sigma_i}{1 - \gamma_i}}, & \gamma_k \neq 0, \\ \frac{\theta}{1 + \frac{1}{-\sigma_k \log p}\Pi} + \frac{(1 - \theta)}{1 + \frac{1}{\sigma_k} \sum_{i=1, i \neq k}^{I+1} \frac{\sigma_i}{1 - \gamma_i}}, & \gamma_k = 0, \end{cases} \quad (18)$$

where θ is a real constant satisfying $0 \leq \theta \leq 1$ and

$$\begin{aligned} \Pi &= \sum_{i=1, i \neq k}^{I+1} \left[\frac{\sigma_i}{\gamma_i} (p^{-\gamma_i} - 1) \mathbb{I}_{\{X_i: \gamma_i \neq 0\}}(X_i) - \sigma_i \log p \mathbb{I}_{\{X_i: \gamma_i = 0\}}(X_i) \right], \\ \frac{\sigma_{I+1}}{\gamma_{I+1}(p^{-\gamma_{I+1}} - 1)} &= \min_{i=1, \dots, I} \left\{ \frac{\sigma_i}{\gamma_i(p^{-\gamma_i} - 1)} \right\}, \\ \frac{\sigma_{I+1}}{1 - \gamma_{I+1}} &= \min_{i=1, \dots, I} \left\{ \frac{\sigma_i}{1 - \gamma_i} \right\}, \\ \sigma_{I+1} &= \min_{i=1, \dots, I} \{\sigma_i\}. \end{aligned}$$

Proposition 7. Consider the normalized GP Mean-VaR measure of risk $\bar{R}_2^{GP}(\gamma_k, \sigma_k, \theta; p)$ defined in (18) with $\gamma_k \neq 0$ and let

$$G(\gamma_k) = (1 - \gamma_k) \sum_{i=1, i \neq k}^{I+1} \frac{\sigma_i}{1 - \gamma_i} - \frac{\gamma_k}{p^{-\gamma_k} - 1} \Pi.$$

Then, $\bar{R}_2^{GP}(\gamma_k, \sigma_k, \theta; p)$ increases with σ_k ; increases with γ_k ; increases with θ if $G(\gamma_k) > 0$, decreases with θ if $G(\gamma_k) < 0$, and $\bar{R}_2^{GP}(\gamma_k, \sigma_k, \theta; p)$ is constant in θ if $G(\gamma_k) = 0$. In the case where $\gamma_k = 0$, the normalized GP Mean-VaR measure of risk increases with σ_k , increases with θ if

$$\sum_{i=1, i \neq k}^{I+1} \frac{\sigma_i}{1 - \gamma_i} + \frac{1}{\log p} \Pi > 0, \text{ decreases with } \theta \text{ if } \sum_{i=1, i \neq k}^{I+1} \frac{\sigma_i}{1 - \gamma_i} + \frac{1}{\log p} \Pi < 0, \text{ and } \bar{R}_2^{GP}(\gamma_k, \sigma_k, \theta; p) \text{ is constant in } \theta \text{ if } \sum_{i=1, i \neq k}^{I+1} \frac{\sigma_i}{1 - \gamma_i} + \frac{1}{\log p} \Pi = 0.$$

4. ARAS representation of the risk models

The aim of multi-criteria decision-making (MCDM) methods is to evaluate and select the best option from a set of alternatives or to classify the alternatives providing a ranking, taking into account multiple criteria. In those methods, quantitative criteria values are normalized, which is fundamental for obtaining comparable scales of criteria values. Different weights, objective or subjective, can be assigned to the criteria in order to reflect the relative importance of each criterion when evaluating and comparing the alternatives. Since decisions often involve multiple conflicting criteria, the weights help quantifying the influence each criterion should have in the final decision-making process. There exist different objective weighting procedures for defining the criteria weights (see e.g. Chatterjee and Chakraborty (2024), Hwang et al. (1993)), for example, equal (or mean) weight method, standard deviation method, coefficient of variation method, Criteria Importance Through Intercriteria Correlation method (CRITIC) method.

The Additive Ratio Assessment (ARAS) method is a multi-criteria decision-making method that was developed in Zavadskas and Turskis (2010). In the ARAS method, the degree of utility of each alternative from the optimal solution is determined through a ratio between the sum aggregated solutions for each alternative and the optimal solution. The ARAS method adapted to risk assessment using risk criteria can be described as follows (see Zavadskas and Turskis (2010), and Brito (2025) for the adaptation to risk assessment).

Definition 10 (ARAS method with risk criteria). Consider the alternatives $A_i, i = 1, \dots, I$, and criteria defined by risk measures $RM_j, j = 1, \dots, J$. Let $X = [x_{ij}]_{I \times J}$ be the decision matrix, where x_{ij} is the risk value of the alternative A_i determined with the risk measure RM_j . The ARAS method with risk criteria consists of the following steps:

Step 1 Calculate the normalized risk decision matrix $R_A = [r_{ij}]_{(I+1) \times J}$ using

$$r_{ij} = \frac{x_{ij}}{\sum_{i=1}^{I+1} x_{ij}}, \quad (19)$$

where $x_{I+1,j}$, $j = 1, \dots, J$, are the optimal values for the J risk criteria.

Step 2 Compute the weighted normalized risk decision matrix $W_A = [w_{ij}]_{(I+1) \times J}$ through

$$w_{ij} = w_j r_{ij}, \quad i = 1, \dots, I+1, \quad j = 1, \dots, J, \quad (20)$$

using criteria weights satisfying $\sum_{j=1}^J w_j = 1$.

Step 3 For each alternative calculate the optimality function:

$$S_i = \sum_{j=1}^J w_{ij}, \quad i = 1, \dots, I+1. \quad (21)$$

Step 4 For each alternative, determine the utility degree k_i , given by

$$k_i = \frac{S_{I+1}}{S_i}, \quad i = 1, \dots, I+1, \quad (22)$$

where S_{I+1} is the optimal, ideal best value.

Step 5 Rank the alternatives according to their utility degree, where the alternative with the maximum utility degree is ranked as the best alternative.

We will show that the normalized risk decision models proposed in Sections 2 and 3, the normalized Mean-SD and the normalized Mean-VaR risk models and the corresponding GP risk models, can be identified with the ARAS method when considering two risk criteria (in the present case: mean and standard deviation, mean and VaR) and by defining the trade-off factor as weighting factor determined with a given weighting procedure.

Remark 1. Consider the alternatives A_i with random variables X_i , $i = 1, \dots, I$. The normalized Mean-SD and Mean-VaR risk models correspond to the ARAS method with risk criteria $RM_1 = \text{SD}[X]$ and $RM_1 = \text{VaR}[X; p]$, respectively, and $RM_2 = \mathbb{E}[X]$, by identifying the trade-off factor with the criteria weights $\theta = w_1$ and $1 - \theta = w_2$, where the risk measures in Definition 2 and Definition 4 correspond to the optimality function in Definition 10, i.e.

$$\bar{R}_l(X_i, \theta) = S_i, \quad l = 1, 2.$$

The risk order relation in Definition 5 leads to the same optimal solution and to the same risk decision as using the utility degree in Definition 10:

Given two alternatives A_{i_1} and A_{i_2} , $i_1, i_2 \in \{1, \dots, I\}$, then

$$[\bar{R}_l(X_{i_1}, \theta) < \bar{R}_l(X_{i_2}, \theta) \Rightarrow A_{i_1} < A_{i_2}] \Leftrightarrow [k_{i_1} > k_{i_2} \Rightarrow A_{i_1} < A_{i_2}], \quad (23)$$

for $l = 1, 2$.

In particular, note that considering the risk decision matrix $X = [x_{ij}]_{(I+1) \times 2}$, where $x_{i1} = \text{SD}[X_i]$, $x_{i2} = \mathbb{E}[X_i]$, for $i = 1, \dots, I$, and

$$x_{I+1,1} = \text{SD}[X_{I+1}] = \min_{X_i} \{\text{SD}[X_i]\},$$

$$x_{I+1,2} = \mathbb{E}[X_{I+1}] = \min_{X_i} \{\mathbb{E}[X_i]\},$$

the optimality function S_i in Step 3 of Definition 10

$$S_i = w_{i1} + w_{i2} = w_1 r_{i1} + w_2 r_{i2}, \quad (24)$$

with

$$r_{i1} = \frac{\text{SD}[X_i]}{\sum_{n=1}^{I+1} \text{SD}[X_n]}, \quad r_{i2} = \frac{\mathbb{E}[X_i]}{\sum_{n=1}^{I+1} \mathbb{E}[X_n]}, \quad (25)$$

corresponds to the normalized Mean-SD risk measure

$$\bar{R}_1(X_i, \theta) = \theta \text{SD}[X_i] + (1 - \theta) \mathbb{E}[X_i], \quad (26)$$

where $\theta = w_1$ and $1 - \theta = 1 - w_1 = w_2$. Similarly, the optimality function S_i can be identified with the normalized Mean-VaR risk measure by employing the criterion $RM_1 = \text{VaR}[X; p]$ instead of $RM_1 = \text{SD}[X]$.

Regarding the determination of the optimal trade-off parameter θ in the here considered risk models, since this parameter corresponds to the weight attributed to the first criterion (see Remark 1: $\theta = w_1$, $1 - \theta = w_2$), and is related with the criteria weight or criteria importance determination in MCDM methods, the weights, and thus the trade-off parameter, can be obtained using an objective or a subjective weighting method, as mentioned before. The subjective weighting methods determine the weights based on the decision-makers' preferences. The objective weighting methods determine the weights based on the available data sample and variation of data. For the present application, an objective weighting method is more adequate (since a subjective judgement to find

the optimal combination of the considered risk measures to obtain the best decision is complex). The standard deviation procedure is one of the objective weighting methods. This procedure assigns higher weights for criteria that differentiate more the risk values of the alternatives, yielding a better impact on the differentiation of alternatives (see e.g. Vavrek (2019), where the use of this method is recommended for the MCDM method TOPSIS). In this approach, the standard deviation procedure will be applied to determine the optimal weight (trade-off parameter). In this context, optimal weight means that the associated risk criterion has a higher standard deviation and therefore a greater ability to differentiate the alternatives. A risk criterion that varies more across the alternatives is treated as more informative and therefore a higher weight will be assigned. Using this procedure, the weights are computed by

$$w_j = \frac{\sigma_j}{\sum_{k=1}^J \sigma_k}, \quad (27)$$

where σ_j is the standard deviation of $(r_{1j}, \dots, r_{Ij})$ (i.e. the standard deviation of the normalized risk values of the j -th risk criterion), $j = 1, \dots, J$. According to this method, the trade-off parameter for the risk models, that will be denoted by $\bar{\theta}$, is defined as

$$\bar{\theta} = \frac{\sigma_1}{\sigma_1 + \sigma_2}, \quad (28)$$

where: for the normalized Mean-SD model, σ_1 is the standard deviation of $(\overline{SD}[X_1], \dots, \overline{SD}[X_J])$; for the normalized Mean-VaR model, σ_1 is the standard deviation of $(\overline{VaR}[X_1; p], \dots, \overline{VaR}[X_J; p])$; σ_2 is the standard deviation of $(\overline{E}[X_1], \dots, \overline{E}[X_J])$.

5. Application to air pollution risk assessment due to PM_{2.5}

The risk models developed in the previous sections will now be applied to assess the risk of air pollution caused by PM_{2.5}, taking into account the monthly nature of the data. The aim is to evaluate the risk of air pollution throughout the year and to identify the months with the highest risk. Monthly data of PM_{2.5} in Antwerp (Belgium) in the period 2015 to 2023, extracted from the European Environment Agency (European Environment Agency, 2025), will be used in the analysis. We exclude the atypical pandemic years 2020–2021.

5.1. GP parameter estimation

The GP distribution was fitted to each month's exceedances by maximum likelihood (ML) within the peaks-over-threshold (POT) framework. The POT approach focuses on modeling only the most extreme values in a dataset by selecting a high threshold and analyzing the exceedances above it. This method allows the tail behavior to be captured using the GP distribution, providing a flexible and efficient framework for extreme-value analysis. The ML estimators of the GP parameters σ and $\gamma \neq 0$, with reparametrization $\tau = \gamma/\sigma$, satisfy the following equations:

$$\begin{cases} \frac{1}{\hat{\tau}} - \left(\frac{1}{\hat{\gamma}} + 1 \right) \frac{1}{N_\mu} \sum_{i=1}^{N_\mu} \frac{X_i}{1 + \hat{\tau} X_i} = 0 \\ \hat{\gamma} = \frac{1}{N_\mu} \sum_{i=1}^{N_\mu} \log(1 + \hat{\tau} X_i), \end{cases} \quad (29)$$

where $X_1, \dots, X_{N_\mu}$ are the N_μ excesses of the chosen high threshold μ . If $\gamma = 0$, then $\hat{\sigma} = \bar{X}$. The selection of an appropriate high threshold μ is a key issue in extreme value theory, with several approaches available in the literature (see, e.g., Beirlant et al. (2004) and references therein). A low threshold can bias the fit, while a high one increases variance, making threshold choice a balance between bias and variance in tail estimation. The approach adopted in this study relies on the GP coefficient of variation (CV), originally proposed by Castillo and Padilla (2016) and further examined by Moreira and Ferreira (2025). This graphical and statistical procedure exploits the property that, for sufficiently high thresholds, the CV of the excesses remains constant when the data follow a GP distribution. By computing the CV across a sequence of increasing thresholds and testing its constancy via a parametric-bootstrap approach, one identifies the smallest threshold at which the null hypothesis of constant CV cannot be rejected, thus yielding a data-driven, GP-consistent cut-off. This method is implemented in the `ercv` package (Castillo et al., 2019) for R (Core, 2020). In Fig. 1 we plot the empirical complementary distribution function, i.e., the empirical $1 - H$ and the respective adjusted model GP. The closeness between both curves suggests an adequate fit of the model. The vertical lines correspond to the threshold μ selected using the method described above. Table 1 reports the estimates of the GP distribution parameters together with the 95% confidence intervals constructed using the normal approximation, $\hat{\omega}_j \pm 1.96 SE(\hat{\omega}_j)$, $j = 1, 2$, where $\hat{\omega} = (\hat{\sigma}, \hat{\gamma})$. The standard errors (SE) are obtained from the inverse of the observed information matrix, computed by numerically evaluating the Hessian of the log-likelihood at the estimates. When the GP model is appropriate, the confidence intervals for the shape and scale parameters generally cover the estimates obtained over nearby plausible thresholds, indicating stability of the tail behaviour. In the next two sections, we will also analyse the effect of GP modelling uncertainty based on bootstrap resampling of the excesses above the selected threshold in each month.

Table 2 presents the risk measures associated with the models, namely the mean, standard deviation and value at risk for $p = 0.01$ and $p = 0.05$ derived, respectively, from (12)–(14), by replacing σ and γ with their respective estimates.

The optimal values of the risk measures (i.e. the lowest values) are highlighted in bold in Table 2, which belong all to July. These values are considered in the normalization process for deriving the normalized GP risk measures.

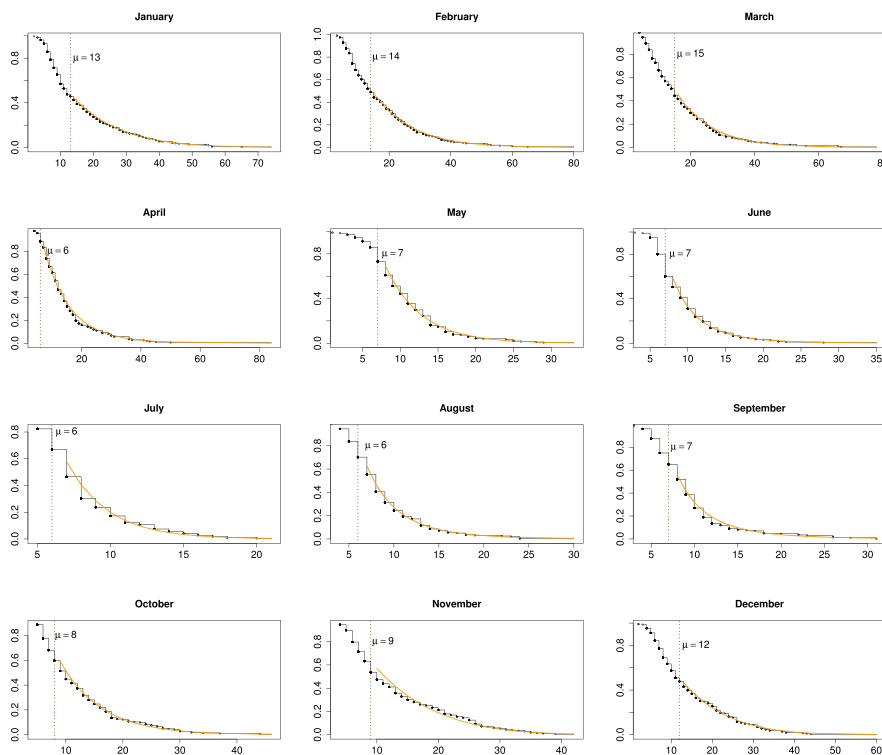


Fig. 1. Empirical complementary distribution function and the respective adjusted model GP (orange curves), from January to December. The vertical dashed lines correspond to the threshold μ .

Table 1

Estimates of GP parameters $\hat{\gamma}$ and $\hat{\sigma}$ and respective 95% confidence intervals, $(\hat{\gamma}_{inf}, \hat{\gamma}_{sup})$ and $(\hat{\sigma}_{inf}, \hat{\sigma}_{sup})$.

Month	$\hat{\gamma}$	$(\hat{\gamma}_{inf}, \hat{\gamma}_{sup})$	$\hat{\sigma}$	$(\hat{\sigma}_{inf}, \hat{\sigma}_{sup})$
Jan	-0.0973	(-0.2819, 0.0872)	14.5742	(10.3446, 18.8038)
Feb	-0.0266	(-0.2032, 0.1501)	12.6342	(8.8490, 16.4195)
Mar	-0.0017	(-0.1879, 0.1845)	10.7264	(7.2908, 14.1620)
Apr	0.0514	(-0.0619, 0.1648)	8.6523	(6.8384, 10.4662)
May	-0.0050	(-0.1305, 0.1204)	4.7854	(3.4861, 6.0848)
Jun	0.1429	(0.0151, 0.2706)	3.1541	(1.9783, 4.3298)
Jul	-0.0019	(-0.1496, 0.1458)	2.8192	(1.9477, 3.6908)
Aug	0.1200	(-0.0237, 0.2638)	3.3261	(2.2673, 4.3848)
Sep	0.1756	(0.0134, 0.3378)	3.2612	(2.2221, 4.3004)
Oct	-0.0019	(-0.1599, 0.1561)	7.1509	(5.0371, 9.2647)
Nov	-0.1606	(-0.3131, -0.0081)	9.9854	(6.7708, 13.2001)
Dec	-0.1460	(-0.2756, -0.0165)	11.3116	(8.4443, 14.1789)

Table 2

Risk components for $\hat{\gamma}$ and $\hat{\sigma}$.

Month	SD[X]	$\mathbb{E}[X]$	VaR[X;0.01]	VaR[X;0.05]
Jan	12.1515	13.2816	54.0925	37.8717
Feb	11.9925	12.3071	54.7629	36.3811
Mar	10.6900	10.7082	49.2038	32.0516
Apr	9.6305	9.1216	44.9613	28.0241
May	4.7375	4.7614	21.7834	14.2280
Jun	4.3541	3.6798	20.5491	11.7930
Jul	2.8087	2.8140	12.9271	8.4220
Aug	4.3360	3.7798	20.4519	11.9911
Sep	4.9111	3.9559	23.1200	12.8560
Oct	7.1374	7.1241	32.7887	21.3618
Nov	7.4853	8.6038	32.4998	23.7459
Dec	8.6832	9.8702	37.9214	27.4464

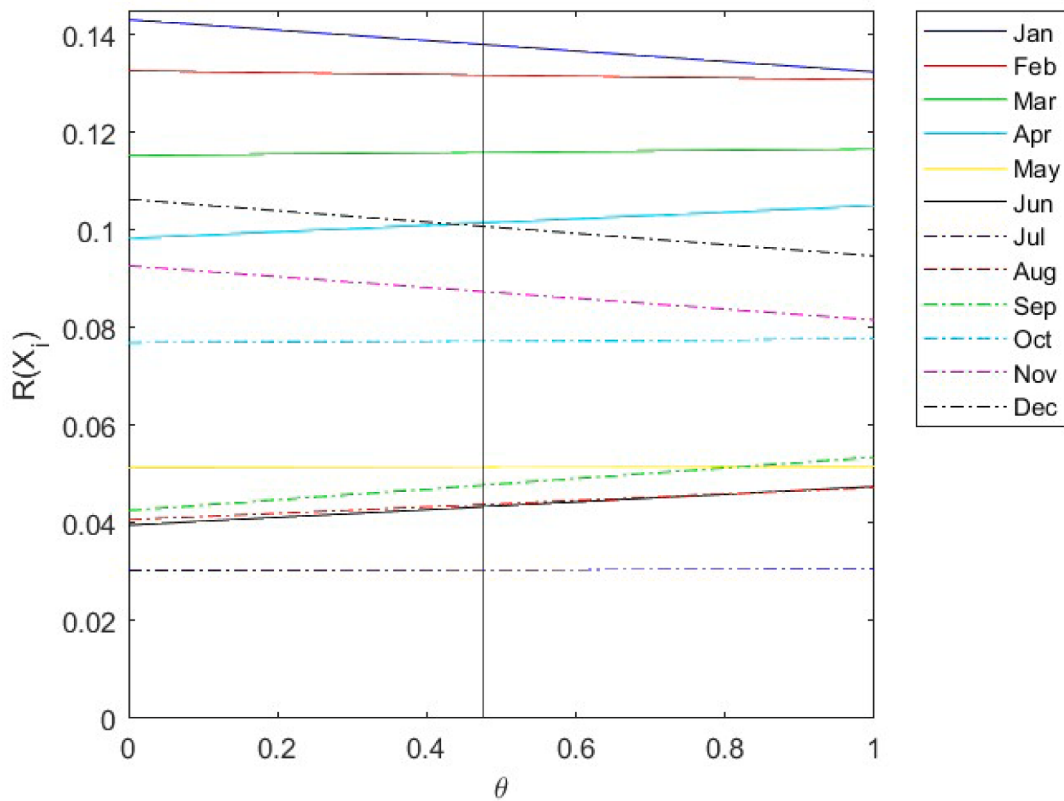


Fig. 2. Normalized Mean-SD risks, where the vertical line corresponds to $\theta = 0.4752$ obtained in (30).

Note that, statistically, the reason July has the lowest risk components is that it attained the lowest scale value, $\hat{\sigma}$ (cf. Table 1). Although the expressions of the risk components also depend on the shape parameter and five months exhibit a lower $\hat{\gamma}$ value than July, the combined effect of both parameters, driven by the considerably lower $\hat{\sigma}$, implies that July is the lowest risk month. From a physical point of view, the lower fine particulate matter concentrations in July can be explained by the fact that heating related emissions are minimal during this hot month, and that summer conditions with higher temperatures enhance atmospheric mixing and dispersion, which in turn promote the dilution of these pollutants.

5.2. Normalized GP Mean-SD risk model

Computing the normalized risk measures yields the results in Table B.7 in Appendix B, which are also plotted in Fig. B.4 (see Appendix B), where a pattern is visible: the colder months show higher values, which gradually decrease towards the warmer months, culminating in the lowest values in the month of July.

The normalized Mean-SD risk measures for the 12 months are plotted in Fig. 2 for $\theta \in [0, 1]$. One can observe that for some months, e.g. January or December, the risks decrease with θ . According to Proposition 5, for the alternatives with positive factor $F(\gamma_k)$, risk increases with θ and for those with negative factor, risk decreases with θ . Calculating $F(\gamma_k)$ (see the results in Table B.8 in Appendix B), one can confirm that for January, February, November and December the risk decreases with θ , whereas for the other months the risk increases with θ .

Fig. 2 shows that, in general, the ordering of the alternatives is relatively stable with respect to the parameter θ . In fact, the variation of θ only changes:

- the risk orders of April and December, where for $\theta = 0$ December has a higher risk than April, however, since April has an increasing risk and December a decreasing risk, by increasing θ , April surpasses December in terms of risk at $\theta = 0.4372$;
- the risk orders of May and September, where for $\theta = 0$ May has a higher risk than September, however a change in the risk orders occurs at $\theta = 0.8286$, so that e.g. for $\theta = 1$ the risk in September is higher than in May.

Thus, according to the normalized Mean-SD risk model the months are ranked as presented in Table 3 for interval ranges $I_\theta \subset [0, 1]$.

From the results in Table 3, one can conclude that in Antwerp the cold period from November to April represents a higher risk for air pollution due to $PM_{2.5}$ than the milder and warmer period ranging from May to October, where January is ranked as riskiest month, followed by February and March, and, on the contrary, July achieves the lowest risk and June and August, respectively, the second and third lowest risks.

Table 3Mean-SD risk ordering of months for interval ranges $I_\theta \subset [0, 1]$.

I_θ	Risk ordering
[0, 0.4372)	Jan > Feb > Mar > Dec > Apr > Nov > Oct > May > Sep > Aug > Jun > Jul
[0.4372, 0.8286)	Jan > Feb > Mar > Apr > Dec > Nov > Oct > May > Sep > Aug > Jun > Jul
[0.8286, 1]	Jan > Feb > Mar > Apr > Dec > Nov > Oct > Sep > May > Aug > Jun > Jul

Table 4

Optimality values S_i , utility degrees k_i , rank orders using ARAS, with mean and standard deviation, and standard deviation weighting procedure (the rank order 1 corresponds to the best classified, i.e. least risky, alternative and the rank order 12 to the worst classified alternative, i.e. the one with the highest risk) and Bootstrap resampling results ($SD(S_i)$, MRP, MR, $SD(R)$, MedR).

month	S_i	k_i	rank	$SD(S_i)$	MRP	MR	$SD(R)$	MedR
Jan	0.1380	0.2207	12	0.0114	10.4289	12	3.1292	11
Feb	0.1317	0.2313	11	0.0124	10.1238	11	2.7559	11
Mar	0.1159	0.2628	10	0.0143	9.5889	10	2.3459	10
Apr	0.1015	0.3002	9	0.0099	8.1400	8	1.6090	8
May	0.0515	0.5919	5	0.0042	4.2697	5	0.8143	4
Jun	0.0434	0.7025	2	0.0147	3.6841	3	1.7045	3
Jul	0.0305	1.0000	1	0.0041	1.6324	1	1.7820	1
Aug	0.0438	0.6949	3	0.0060	3.2905	2	1.7770	3
Sep	0.0478	0.6371	4	0.0077	4.1780	4	1.8773	4
Oct	0.0773	0.3943	6	0.0071	6.4673	6	1.2547	6
Nov	0.0874	0.3484	7	0.0058	7.4452	7	1.3030	7
Dec	0.1008	0.3022	8	0.0085	8.7512	9	1.3328	9

The previous results show that, the trade-off parameter in the present application has, in general, a low influence on the final ranking result obtained with the normalized Mean-SD risk model. By using the normalized Mean ($\theta = 0$), or the normalized SD ($\theta = 1$) risk model, the differences in the final rankings are changes in the positions of four alternatives with intermediate risk levels (cf. Table 3: April-December, May-September). Next, the ARAS method will be considered, along with the standard deviation weighting procedure for the determination of the trade-off parameter θ .

ARAS with mean and standard deviation criteria

Since the ranking obtained with the ARAS method for any weighting method using the criteria mean and standard deviation must equal one of the rankings in Table 3, one can confirm the ARAS output result by selecting a specific weighting strategy. As explained in Section 4, here, the standard deviation weighting procedure (28), will be applied, yielding

$$\bar{\theta} = w_1 = 0.4752, \quad 1 - \bar{\theta} = w_2 = 0.5248, \quad (30)$$

with output results: S_i , k_i and rank, listed in Table 4. Thus, one can confirm that the ranking obtained with the trade-off parameter $\bar{\theta} = 0.4752$ coincides in fact with the ranking obtained for $\theta \in [0.4372, 0.8286]$ in Table 3.

Since Table 3 presents results for a continuous variation of the parameter θ within the interval $[0, 1]$, the sensitivity analysis conducted on the weighting factor demonstrates the method's strong robustness, as only two rank reversals are observed.

Uncertainty analysis for the normalized Mean-SD risk model

In order to quantify and analyse uncertainty using the resampling-based approach, we generate 10,000 bootstrap replicates from the monthly excesses and re-estimate the Generalized Pareto model for each sample. For each bootstrap replicate, the normalized risk S_i , $i = 1, \dots, 12$, as defined in (24), is computed according to the Mean-SD model. Then, the median ranking (MedR) of the 10,000 Bootstrap rankings is determined (see Table 4). One can observe that the median rank produces ties (January and February; May and September; June and August). To ensure a complete and fully ordered ranking, the mean of the Bootstrap ranks (MR), obtained from the mean ranking positions (MRP) of the 10,000 Bootstrap rankings (see Table 4), is employed as a tie-breaking criterion. This approach preserves the robustness of the median while allowing for an unambiguous ranking. The resulting final ranking (obtained by breaking the ties in MedR using MR) coincides with the mean ranking. Therefore, the final Bootstrap ranking is given by MR.

Comparing the final Bootstrap ranking (MR) with the ranking previously obtained using the proposed methodology, one can observe that they are very similar; only two rank reversals occur, namely between April and December, and between June and August, each differing by just one rank and now swapping their relative risk positions. Computing the Spearman's rank correlation coefficient and Kendall's tau coefficient for examining the similarity of both rankings, the results reveal a strong agreement of these rankings, since the correlation coefficients belong to the interval $[0.93, 1]$ and the p -values are lower than 0.01, indicating a significant correlation between the rankings.

The uncertainty of the monthly rankings was assessed using bootstrap-based standard deviations of both the estimated scores S_i and the ranks, denoted by $SD(S_i)$ and $SD(R)$, respectively (see Table 4). The variability of S_i is generally low across all months, indicating stable estimation of the underlying scores, although slightly higher values are observed for March (0.0143) and June (0.0147). In contrast, May and July exhibit the lowest variability, suggesting more precise estimation. The variability of the rankings,

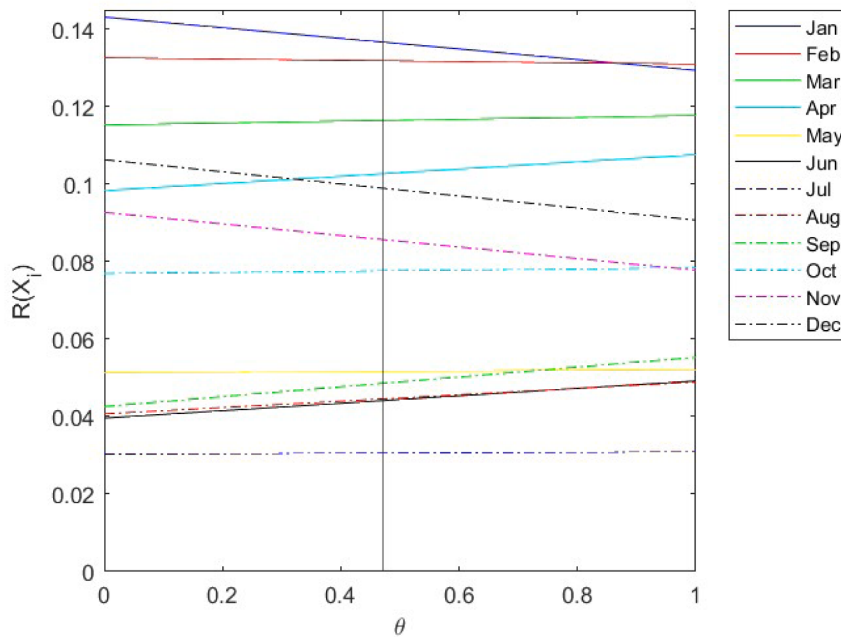


Fig. 3. Normalized Mean-VaR risks for $p = 0.01$, where the vertical line corresponds to $\theta = 0.4708$ obtained in (31).

Table 5

Mean-VaR risk ordering ($p = 0.01$) of months for interval ranges $I_\theta \subset [0, 1]$.

I_θ	Risk ordering
$[0, 0.3213)$	Jan > Feb > Mar > Dec > Apr > Nov > Oct > May > Sep > Aug > Jun > Jul
$[0.3213, 0.7311)$	Jan > Feb > Mar > Apr > Dec > Nov > Oct > May > Sep > Aug > Jun > Jul
$[0.7311, 0.8678)$	Jan > Feb > Mar > Apr > Dec > Nov > Oct > Sep > May > Jun > Aug > Jul
$[0.8678, 0.9634)$	Feb > Jan > Mar > Apr > Dec > Nov > Oct > Sep > May > Jun > Aug > Jul
$[0.9634, 1]$	Feb > Jan > Mar > Apr > Dec > Oct > Nov > Sep > May > Jun > Aug > Jul

measured by $SD(R)$, reveals more pronounced differences across months. In particular, January (3.1292), February (2.7559), and March (2.3459) display the highest ranking uncertainty, indicating substantial sensitivity to resampling. Conversely, May (0.8143) shows the lowest variability, followed by October and November, suggesting more stable ranking positions. Overall, these results suggest that ranking uncertainty is primarily driven by the proximity of the underlying scores, with months exhibiting similar S_i values showing greater instability in their relative ordering.

One can conclude that although the parameter uncertainty causes slight variations in the monthly risk rankings, there is an overall strong agreement among the rankings, suggesting that the ranking obtained for $\hat{\gamma}$ and $\hat{\sigma}$ can be considered representative of the rankings produced from the Bootstrap resampling procedure.

5.3. Normalized GP Mean-VaR risk model

The normalized GP Mean-VaR risk model will be studied for the significance levels $p = 0.01$ and $p = 0.05$, where the results for $p = 0.01$ are presented in this section and the results for $p = 0.05$ are given in Appendix C.

Fig. 3 contains the graphical representations of the normalized Mean-VaR risk measures, defined with the normalized components $\overline{\text{VaR}}[X; 0.01]$ and $\overline{\mathbb{E}}[X]$ listed in Table B.7 (see Appendix B), for the twelve months. By varying the parameter θ , changes in risk ordering occur between: December and April, May and September, January and February, October and November (see Table 5).

Calculating the factors $G(\gamma_k)$ of Proposition 7, one obtains the results in Table B.9 (see Appendix B). Since $G(\gamma_k) < 0$ for January, February, November and December, according to Proposition 7, this implies that the risk measure decreases with respect to θ for these months. Conversely, for the remaining months, the risk measure increases with θ . This behaviour of the risk lines is illustrated in Fig. 3.

Analysing the results in Table 5, one can conclude that also the normalized Mean-VaR risk model classifies the cold period from November to April as riskiest period for air pollution due to $\text{PM}_{2.5}$, where January is classified as riskiest month for $\theta \in [0, 0.8678)$ and February as second riskiest month. For $\theta \in [0.8678, 1]$ these months reverse their ranks. July has the lowest risk for $\theta \in [0, 1]$, and June and August are classified as presenting the second lowest risk for $\theta \in [0, 0.7311)$ and for $\theta \in [0.7311, 1]$, respectively.

Table 6

Optimality values S_i , utility degrees k_i , rank orders using ARAS, with mean and VaR with $p = 0.01$, and standard deviation weighting procedure (the rank order 1 corresponds to the best classified, i.e. least risky, alternative and the rank order 12 to the worst classified alternative, i.e. the one with the highest risk) and Bootstrap resampling results ($SD(S_i)$, MRP, MR, $SD(R)$, MedR).

Month	S_i	k_i	rank	$SD(S_i)$	MRP	MR	$SD(R)$	MedR
Jan	0.1366	0.2240	12	0.0271	10.7795	12	1.0504	11
Feb	0.1318	0.2321	11	0.0282	10.7721	11	1.1193	11
Mar	0.1165	0.2628	10	0.0315	10.6509	10	1.1734	11
Apr	0.1026	0.2981	9	0.0244	9.1859	9	1.1873	9
May	0.0517	0.5922	5	0.0107	3.4698	3	0.9449	3
Jun	0.0441	0.6936	2	0.1685	4.2449	5	2.5157	4
Jul	0.0306	1.0000	1	0.0080	1.1320	1	0.5134	1
Aug	0.0446	0.6864	3	0.0117	2.9051	2	1.1116	3
Sep	0.0486	0.6298	4	0.0139	4.0380	4	1.2021	4
Oct	0.0776	0.3943	6	0.0166	6.6114	6	1.0112	7
Nov	0.0857	0.3573	7	0.0162	6.6451	7	0.9537	7
Dec	0.0990	0.3092	8	0.0202	7.5653	8	1.2313	8

ARAS with mean and VaR with $p = 0.01$ criteria

Applying the ARAS method using the standard deviation weighting procedure (28), yields

$$\tilde{\theta} = w_1 = 0.4708, \quad 1 - \tilde{\theta} = w_2 = 0.5292, \quad (31)$$

with output results listed in Table 6. As one can observe, the ARAS ranking coincides with the risk ordering in Table 5 for the trade-off parameter $\theta = 0.4708 \in [0.3213, 0.7311]$.

Regarding the sensitivity analysis based on the variation of the weighting factor (see Table 5), five rank reversals occur, indicating that the VaR with $p = 0.01$ introduces instability in the final ranking.

Uncertainty analysis for the normalized Mean-VaR risk model with $p = 0.01$

Using the resampling-based approach for uncertainty analysis, previously described and applied to the Mean-SD risk model, we computed the normalized risk S_i , $i = 1, \dots, 12$, as defined in (24), for each of the 10,000 bootstrap replicates, now considering the Mean-VaR model with $p = 0.01$. Table 6 contains the median rank (MedR) of the 10,000 Bootstrap rankings and also the mean ranking positions (MRP) with the associated mean rank (MR), that will be used to solve the ties generated by the median rank. One can notice that in this case MedR leads to the following tied months: January and February and March, May and August; June and September; October and November. Breaking the ties using the mean ranking leads exactly to the ranking provided by MR, so that this will be the resulting final Bootstrap ranking (obtained from MedR using MR to break ties).

Comparing both rankings, the rank obtained from the proposed methodology with weights given in (31) and the final Bootstrap ranking (MR), there are three months with differing positions: May, June and August, so that, overall, the rankings are very similar. The Spearman and Kendall correlation analyses show that the correlation coefficients lie in the interval $(0.87, 1]$, and all p -values are below 0.01, indicating a strong and statistically significant agreement between the rankings.

The uncertainty in the monthly rankings was quantified using bootstrap-based standard deviations of both the estimated scores, $SD(S_i)$, and the ranks, $SD(R)$, presented in Table 6. In contrast to the Mean-SD risk model, the variability of S_i is more heterogeneous, with particularly high variability observed for June ($SD(S_i) = 0.1685$), indicating substantial instability in its underlying score, while most other months exhibit moderate to low variability. The bootstrap variability of the rankings, $SD(R)$, is generally lower than in the previous model, suggesting improved overall ranking stability. Nevertheless, June again stands out with the highest ranking uncertainty ($SD(R) = 2.52$), indicating a strong sensitivity to resampling and a high likelihood of rank changes. Moderate variability is observed for March, April, September, and December, whereas months such as July and May display relatively stable rankings. Overall, we conclude that ranking uncertainty is mainly influenced by both the variability and closeness of the underlying scores, with months having similar S_i values displaying greater instability in their relative positions, whereas months with more distinct S_i values tend to exhibit more stable rankings.

Thus, although uncertainty in modelling introduces minor variations, the rankings remain highly consistent, and the ranking based on $\hat{\gamma}$ and $\hat{\sigma}$ can also in this case be regarded as representative.

5.4. General analysis of the results and general remarks

Analysing and comparing the previous results obtained with the three normalized risk models (Mean-SD, Mean-VaR with $p = 0.01$, Mean-VaR with $p = 0.05$), one can conclude the following.

Regarding the sensitivity analysis based on the variation of the trade-off factor, the VaR with $p = 0.01$ introduces a greater instability in the final ranking of the months compared to the SD and to the VaR with $p = 0.05$. Specifically, five rank reversals occur in the Mean-VaR $p = 0.01$ model, two in the Mean-SD model and only one in the Mean-VaR $p = 0.05$ model. The patterns of the Mean-VaR $p = 0.05$ risk lines, with relatively more constant trends than those of the other two risk models, suggest a quite stable risk

classification of the months. Thus, the Mean-VaR $p = 0.05$ risk model is the most stable one with respect to variations in the trade-off parameter.

The bootstrap resampling approach to uncertainty analysis offers important insights into the robustness of the monthly risk rankings. The bootstrap standard deviations of the ranks, $SD(R)$, and of the normalized risk estimates, $SD(S_i)$, reveal distinct patterns across the three models and across months. Overall, $SD(R)$ is substantially larger under the Mean-SD model, particularly in the first months (January, February, March), indicating greater instability in rank positions. In contrast, both Mean-VaR models yield markedly lower values, with the Mean-VaR $p = 0.05$ model generally exhibiting the smallest variability in ranks across most months. However, an exception occurs in June, where the Mean-VaR models show increased variability, exceeding that of the Mean-SD model, suggesting reduced ranking stability under tail-risk measures for this month. Regarding $SD(S_i)$, the Mean-SD model consistently presents the lowest variability across all months. In contrast, both Mean-VaR models exhibit higher dispersion, reflecting increased uncertainty in the estimated risks, especially for extreme quantiles. This effect is particularly pronounced in June. In summary, the Mean-SD model is associated with higher variability in ranking positions but lower uncertainty in risk estimates, whereas the Mean-VaR models provide more stable rankings overall but at the cost of increased variability in the estimated risks, particularly in months with higher tail sensitivity such as June. Despite minor changes in risk ranking positions of some months, the ranking analyses of the Bootstrap rankings, considering the median ranks together with the mean ranks, revealed that the risk rankings obtained for the computed parameter estimates can be considered representative.

The fact that the rankings obtained with the proposed risk models are not substantially different overall indicates that a stable final decision and risk classification of the months can be established. Although in the given application, the ranking appears quite similar to those that would have been obtained with e.g. the simple $\text{VaR}[X; 0.05]$ (see Table 2), the advantage of the proposed risk models lies in their ability to provide a more flexible and diverse risk classification by taking into account two different risk criteria and allowing different weight assignments through the trade-off factor.

In future research, it is of interest to study the risk models in greater depth (particularly by examining how variations in the trade-off factor influence risk) across different datasets and application domains, such as e.g. temperature data in environmental studies or market data in finance.

6. Conclusions

In this paper, an extreme risk modeling approach has been proposed that uses two-component risk measures to assess and classify the risk of alternatives. The new risk measures combine two risk components using a trade-off factor: mean and standard deviation, and mean and value at risk. In order to take into account the features of extreme events, the risk models have been developed for the generalized Pareto distribution and depend on the shape and tail parameters. Properties related with the perception of risk have been analysed and conclusions have been drawn based on the risk model parameters. It is shown that the risk models can be identified with the multi-criteria decision making method: additive ratio assessment (ARAS), when specifying the two considered risk criteria. With the trade-off factor defined to vary continuously within the interval $[0, 1]$, the risk models cover the full range of outputs generated by the ARAS method for different criteria weighting procedures.

The risk models have been applied to air pollution risk assessment due to particulate matter $\text{PM}_{2.5}$ with the aim to identify the most polluted periods in the year and to classify the months in terms of risk. Monthly parameter concentrations of $\text{PM}_{2.5}$ in Antwerp using data from 2015–2023 (excluding the atypical pandemic years 2020–2021) have been analysed. Based on the monthly data samples, the parameters of the underlying GP distributions have been estimated. The Mean-SD and Mean-VaR (with significance levels 1% and 5%) risk models have been defined for the twelve months and their risks have been compared for variations in the trade-off factor. All three risk models have classified the cold period ranging from November to April as the riskiest period for air pollution due to $\text{PM}_{2.5}$ in Antwerp, where January has been identified as the riskiest month followed by February. On the contrary the hot period, with July classified as less riskiest month, has represented the lowest risk for air pollution due to $\text{PM}_{2.5}$. Considering the sensitivity analysis based on the trade-off parameter variation (which equals the weight factor in the ARAS method), one can conclude that the proposed models are very stable for the risk assessment of the alternatives, where only the Mean-VaR risk model with $p = 0.01$ has exhibited more rank reversals (5 alternative pairs changed their ranking positions). In particular, results for the ARAS method with the trade-off factor determined by the standard deviation weighting procedure have been obtained. In order to quantify and analyse the uncertainty in the risk estimation for the normalized risk models, a bootstrap approach has been employed, involving resampling of the excesses above the selected threshold in each month followed by refitting a GP model. Overall, the results indicated that the bootstrap-based rankings were generally close to the rankings obtained from the computed parameter estimates. This analysis provides a comprehensive view of uncertainty, showing that while the bootstrap identifies months sensitive to sampling variability (particularly June), the overall ranking remains largely robust.

The proposed risk assessment methodology, based on the Generalized Pareto distribution, allows for the modeling of extreme values and, consequently, accounts for rare extreme events that may occur in worst-case scenarios. This new risk modeling approach is particularly relevant for environmental risk assessment, where quantifying extreme impacts is essential for effective planning of public health and environmental policies.

In the future, the models are intended to be applied to the risk assessment of air pollution in different locations, including cities in both cold and hot climate regions, and considering different air pollutants. The proposed methodology and the results provide the basis for further developments and improvements of the models, such as exploring alternative parameter estimation procedures, investigating combinations of other different risk measures or incorporating additional risk measures. The methodology will also be extended to the Generalized Extreme Value distribution.

Acknowledgments

This research was partially financed by Portuguese Funds through FCT (“Fundação para a Ciência e a Tecnologia”) within the Project UID/00013/2025 (<https://doi.org/10.54499/UID/00013/2025>). The authors thank the reviewers for helpful comments and suggestions.

Appendix A. Proofs of Propositions in Section 2 and in Section 3

A.1. Proof of Proposition 2

(a) Consider the normalized Mean-SD risk measure $\bar{R}_1$.

1. Since $SD[X] = SD[X + \alpha]$ and $\mathbb{E}[X] < \mathbb{E}[X + \alpha] = \mathbb{E}[X] + \alpha$ for $\alpha > 0$ and $\mathbb{E}[X] > \mathbb{E}[X + \alpha] = \mathbb{E}[X] + \alpha$ for $\alpha < 0$, one has, for $\alpha > 0$:

$$\bar{R}_1(X, \theta) = \frac{\theta}{3} + (1 - \theta) \frac{\mathbb{E}[X]}{3\mathbb{E}[X] + \alpha}, \quad (\text{A.1})$$

$$\bar{R}_1(X + \alpha, \theta) = \frac{\theta}{3} + (1 - \theta) \frac{\mathbb{E}[X] + \alpha}{3\mathbb{E}[X] + \alpha}, \quad (\text{A.2})$$

and, for $\alpha < 0$:

$$\bar{R}_1(X, \theta) = \frac{\theta}{3} + (1 - \theta) \frac{\mathbb{E}[X]}{3\mathbb{E}[X] + 2\alpha}, \quad (\text{A.3})$$

$$\bar{R}_1(X + \alpha, \theta) = \frac{\theta}{3} + (1 - \theta) \frac{\mathbb{E}[X] + \alpha}{3\mathbb{E}[X] + 2\alpha}, \quad (\text{A.4})$$

so that $\bar{R}_1(X + \alpha, \theta) > \bar{R}_1(X, \theta)$, if $\alpha > 0$, and $\bar{R}_1(X + \alpha, \theta) < \bar{R}_1(X, \theta)$, if $\alpha < 0$.

2. Since $SD[\alpha X] = \alpha SD[X]$ and $\mathbb{E}[\alpha X] = \alpha \mathbb{E}[X]$, then, if $\alpha > 1$, one has

$$\bar{R}_1(X, \theta) = \frac{\theta}{2 + \alpha} + \frac{1 - \theta}{2 + \alpha} = \frac{1}{2 + \alpha} \quad (\text{A.5})$$

$$\bar{R}_1(\alpha X, \theta) = \frac{\alpha\theta}{2 + \alpha} + \frac{\alpha(1 - \theta)}{2 + \alpha} = \frac{\alpha}{2 + \alpha}, \quad (\text{A.6})$$

so that $\bar{R}_1(\alpha X, \theta) > \bar{R}_1(X, \theta)$. If $0 < \alpha < 1$, one has

$$\bar{R}_1(X, \theta) = \frac{\theta}{1 + 2\alpha} + \frac{1 - \theta}{1 + 2\alpha} = \frac{1}{1 + 2\alpha} \quad (\text{A.7})$$

$$\bar{R}_1(\alpha X, \theta) = \frac{\alpha\theta}{1 + 2\alpha} + \frac{\alpha(1 - \theta)}{1 + 2\alpha} = \frac{\alpha}{1 + 2\alpha}, \quad (\text{A.8})$$

implying $\bar{R}_1(\alpha X, \theta) < \bar{R}_1(X, \theta)$.

(b) Consider the normalized Mean-VaR risk measure $\bar{R}_2$.

1. Since $\text{VaR}[X; p] < \text{VaR}[X + \alpha; p] = \text{VaR}[X; p] + \alpha$ and $\mathbb{E}[X] < \mathbb{E}[X + \alpha] = \mathbb{E}[X] + \alpha$ for $\alpha > 0$ and $\text{VaR}[X; p] > \text{VaR}[X + \alpha; p] = \text{VaR}[X; p] + \alpha$ and $\mathbb{E}[X] > \mathbb{E}[X + \alpha] = \mathbb{E}[X] + \alpha$ for $\alpha < 0$, one has, for $\alpha > 0$:

$$\bar{R}_2(X, \theta) = \theta \frac{\text{VaR}[X; p]}{3\text{VaR}[X; p] + \alpha} + (1 - \theta) \frac{\mathbb{E}[X]}{3\mathbb{E}[X] + \alpha}, \quad (\text{A.9})$$

$$\bar{R}_2(X + \alpha, \theta) = \theta \frac{\text{VaR}[X; p] + \alpha}{3\text{VaR}[X; p] + \alpha} + (1 - \theta) \frac{\mathbb{E}[X] + \alpha}{3\mathbb{E}[X] + \alpha}, \quad (\text{A.10})$$

and, for $\alpha < 0$:

$$\bar{R}_2(X, \theta) = \theta \frac{\text{VaR}[X; p]}{3\text{VaR}[X; p] + 2\alpha} + (1 - \theta) \frac{\mathbb{E}[X]}{3\mathbb{E}[X] + 2\alpha}, \quad (\text{A.11})$$

$$\bar{R}_2(X + \alpha, \theta) = \theta \frac{\text{VaR}[X; p] + \alpha}{3\text{VaR}[X; p] + 2\alpha} + (1 - \theta) \frac{\mathbb{E}[X] + \alpha}{3\mathbb{E}[X] + 2\alpha}, \quad (\text{A.12})$$

so that $\bar{R}_2(X + \alpha, \theta) > \bar{R}_2(X, \theta)$, if $\alpha > 0$, and $\bar{R}_2(X + \alpha, \theta) < \bar{R}_2(X, \theta)$, if $\alpha < 0$.

2. Since $\text{VaR}[\alpha X] = \alpha \text{VaR}[X; p]$ and $\mathbb{E}[\alpha X] = \alpha \mathbb{E}[X]$, then, if $\alpha > 1$, one has

$$\bar{R}_2(X, \theta) = \frac{\theta}{2 + \alpha} + \frac{1 - \theta}{2 + \alpha} = \frac{1}{2 + \alpha}, \quad (\text{A.13})$$

$$\bar{R}_2(\alpha X, \theta) = \frac{\alpha\theta}{2 + \alpha} + \frac{\alpha(1 - \theta)}{2 + \alpha} = \frac{\alpha}{2 + \alpha}, \quad (\text{A.14})$$

so that $\bar{R}_2(\alpha X, \theta) > \bar{R}_2(X, \theta)$. If $0 < \alpha < 1$, one has

$$\bar{R}_2(X, \theta) = \frac{\theta}{1 + 2\alpha} + \frac{1 - \theta}{1 + 2\alpha} = \frac{1}{1 + 2\alpha}, \quad (\text{A.15})$$

$$\bar{R}_2(\alpha X, \theta) = \frac{\alpha\theta}{1 + 2\alpha} + \frac{\alpha(1 - \theta)}{1 + 2\alpha} = \frac{\alpha}{1 + 2\alpha}, \quad (\text{A.16})$$

implying $\bar{R}_2(\alpha X, \theta) < \bar{R}_2(X, \theta)$.

□

A.2. Proof of Proposition 3

(a) Consider the normalized Mean-SD risk measure $\bar{R}_1$.

1. Considering $\mathbb{X} = \{X_1, X_2\}$, one has

$$\bar{R}_1(X_1, \theta) < \bar{R}_1(X_2, \theta) \Leftrightarrow \theta \frac{\text{SD}[X_1] - \text{SD}[X_2]}{\sum_{i=1}^{I+1} \text{SD}[X_i]} < (1 - \theta) \frac{\mathbb{E}[X_2] - \mathbb{E}[X_1]}{\sum_{i=1}^{I+1} \mathbb{E}[X_i]}. \quad (\text{A.17})$$

Now, let $\mathbb{X} = \{\alpha X_1, \alpha X_2\}$. Due to the fact that

$$\frac{\text{SD}[\alpha X_i]}{\sum_{i=1}^{I+1} \text{SD}[\alpha X_i]} = \frac{\text{SD}[X_i]}{\sum_{i=1}^{I+1} \text{SD}[X_i]},$$

and

$$\frac{\mathbb{E}[\alpha X_i]}{\sum_{i=1}^{I+1} \mathbb{E}[\alpha X_i]} = \frac{\mathbb{E}[X_i]}{\sum_{i=1}^{I+1} \mathbb{E}[X_i]}$$

for $i = 1, 2$ and $I = 2$, where $\text{SD}[\alpha X_3] = \min\{\text{SD}[\alpha X_1], \text{SD}[\alpha X_2]\} = \alpha \min\{\text{SD}[X_1], \text{SD}[X_2]\}$ and $\mathbb{E}[\alpha X_3] = \min\{\mathbb{E}[\alpha X_1], \mathbb{E}[\alpha X_2]\} = \alpha \min\{\mathbb{E}[X_1], \mathbb{E}[X_2]\}$, it follows that $\bar{R}_1(X_1, \theta) < \bar{R}_1(X_2, \theta)$ if and only if $\bar{R}_1(\alpha X_1, \theta) < \bar{R}_1(\alpha X_2, \theta)$ for $\mathbb{X} = \{\alpha X_1, \alpha X_2\}$.

2. Let $\mathbb{X} = \{\alpha X_1, \beta X_1\}$. Assume that $\alpha < \beta$. It follows that

$$\bar{R}_1(\alpha X_1, \theta) < \bar{R}_1(\beta X_1, \theta) \Leftrightarrow \frac{\theta\alpha}{2\alpha + \beta} + \frac{(1 - \theta)\alpha}{2\alpha + \beta} < \frac{\theta\beta}{2\alpha + \beta} + \frac{(1 - \theta)\beta}{2\alpha + \beta} \Leftrightarrow \frac{\alpha}{2\alpha + \beta} < \frac{\beta}{2\alpha + \beta}.$$

and, considering $\mathbb{X} = \{\alpha X_2, \beta X_2\}$, that

$$\bar{R}_1(\alpha X_2, \theta) < \bar{R}_1(\beta X_2, \theta) \Leftrightarrow \frac{\theta\alpha}{2\alpha + \beta} + \frac{(1 - \theta)\alpha}{2\alpha + \beta} < \frac{\theta\beta}{2\alpha + \beta} + \frac{(1 - \theta)\beta}{2\alpha + \beta} \Leftrightarrow \frac{\alpha}{2\alpha + \beta} < \frac{\beta}{2\alpha + \beta},$$

since $\text{SD}[\alpha X_i] = \frac{\alpha}{2\alpha + \beta}$, $\mathbb{E}[\alpha X_i] = \frac{\alpha}{2\alpha + \beta}$, $\text{SD}[\beta X_i] = \frac{\beta}{2\alpha + \beta}$ and $\mathbb{E}[\beta X_i] = \frac{\beta}{2\alpha + \beta}$ for $i = 1, 2$. Therefore, one concludes that given $\mathbb{X} = \{\alpha X_1, \beta X_1\}$ and $\bar{R}_1(\alpha X_1, \theta) < \bar{R}_1(\beta X_1, \theta)$ is equivalent to $\bar{R}_1(\alpha X_2, \theta) < \bar{R}_1(\beta X_2, \theta)$ for $\mathbb{X} = \{\alpha X_2, \beta X_2\}$. The proof for the case $\beta < \alpha$ is analogous, one just has to interchange α and β in the expressions for the normalized quantities and in the inequalities.

(b) Consider the normalized Mean-VaR risk measure $\bar{R}_2$.

1. Given $\mathbb{X} = \{X_1, X_2\}$, one has

$$\bar{R}_2(X_1, \theta) < \bar{R}_2(X_2, \theta) \Leftrightarrow \theta \frac{\text{VaR}[X_1; p] - \text{VaR}[X_2; p]}{\sum_{i=1}^{I+1} \text{VaR}[X_i; p]} < (1 - \theta) \frac{\mathbb{E}[X_2] - \mathbb{E}[X_1]}{\sum_{i=1}^{I+1} \mathbb{E}[X_i]}.$$

Consider $\mathbb{X} = \{\alpha X_1, \alpha X_2\}$. Since

$$\frac{\text{VaR}[\alpha X_i; p]}{\sum_{i=1}^{I+1} \text{VaR}[\alpha X_i; p]} = \frac{\text{VaR}[X_i; p]}{\sum_{i=1}^{I+1} \text{VaR}[X_i; p]}$$

and

$$\frac{\mathbb{E}[\alpha X_i]}{\sum_{i=1}^{I+1} \mathbb{E}[\alpha X_i]} = \frac{\mathbb{E}[X_i]}{\sum_{i=1}^{I+1} \mathbb{E}[X_i]}$$

for $i = 1, 2$ and $I = 2$, where $\text{VaR}[\alpha X_3; p] = \min\{\text{VaR}[\alpha X_1; p], \text{VaR}[\alpha X_2; p]\} = \alpha \min\{\text{VaR}[X_1; p], \text{VaR}[X_2; p]\}$ and $\mathbb{E}[\alpha X_3] = \min\{\mathbb{E}[\alpha X_1], \mathbb{E}[\alpha X_2]\} = \alpha \min\{\mathbb{E}[X_1], \mathbb{E}[X_2]\}$, it follows that $\bar{R}_2(X_1, \theta) < \bar{R}_2(X_2, \theta)$ if and only if $\bar{R}_2(\alpha X_1, \theta) < \bar{R}_2(\alpha X_2, \theta)$ for $\mathbb{X} = \{\alpha X_1, \alpha X_2\}$.

2. Consider $\mathbb{X} = \{\alpha X_1, \beta X_1\}$. Assume that $\alpha < \beta$. Then, one has

$$\bar{R}_2(\alpha X_1, \theta) < \bar{R}_2(\beta X_1, \theta) \Leftrightarrow \frac{\theta\alpha}{2\alpha + \beta} + \frac{(1 - \theta)\alpha}{2\alpha + \beta} < \frac{\theta\beta}{2\alpha + \beta} + \frac{(1 - \theta)\beta}{2\alpha + \beta} \Leftrightarrow \frac{\alpha}{2\alpha + \beta} < \frac{\beta}{2\alpha + \beta}$$

and, considering $\mathbb{X} = \{\alpha X_2, \beta X_2\}$, it turns out that

$$\bar{R}_2(\alpha X_2, \theta) < \bar{R}_2(\beta X_2, \theta) \Leftrightarrow \frac{\theta\alpha}{2\alpha + \beta} + \frac{(1 - \theta)\alpha}{2\alpha + \beta} < \frac{\theta\beta}{2\alpha + \beta} + \frac{(1 - \theta)\beta}{2\alpha + \beta} \Leftrightarrow \frac{\alpha}{2\alpha + \beta} < \frac{\beta}{2\alpha + \beta},$$

since $\text{VaR}[\alpha X_i; p] = \frac{\alpha}{2\alpha + \beta}$, $\mathbb{E}[\alpha X_i] = \frac{\alpha}{2\alpha + \beta}$, $\text{VaR}[\beta X_i; p] = \frac{\beta}{2\alpha + \beta}$ and $\mathbb{E}[\beta X_i] = \frac{\beta}{2\alpha + \beta}$ for $i = 1, 2$. Therefore, one concludes that, given $\mathbb{X} = \{\alpha X_1, \beta X_1\}$, $\bar{R}_2(\alpha X_1, \theta) < \bar{R}_2(\beta X_1, \theta)$ is equivalent to $\bar{R}_2(\alpha X_2, \theta) < \bar{R}_2(\beta X_2, \theta)$ for $\mathbb{X} = \{\alpha X_2, \beta X_2\}$. The proof for $\beta < \alpha$ is analogous, one just has to interchange α and β in the normalized quantities and in the inequalities.

□

A.3. Proof of Proposition 4

From the definition of $R_1^{\text{GP}}(\gamma, \sigma, \theta)$, it is evident that the risk increases with σ and with γ . Rewriting $R_1^{\text{GP}}(\gamma, \sigma, \theta)$ as

$$R_1^{\text{GP}}(\gamma, \sigma, \theta) = \theta \left[\frac{\sigma(1 - \sqrt{1 - 2\gamma})}{(1 - \gamma)\sqrt{1 - 2\gamma}} \right] + \frac{\sigma}{1 - \gamma}, \quad (\text{A.18})$$

one can conclude that $R_1^{\text{GP}}(\gamma, \sigma, \theta)$ increases with θ , whenever

$$\sigma(1 - \sqrt{1 - 2\gamma}) > 0 \Leftrightarrow \gamma > 0, \quad (\text{A.19})$$

decreases with θ , whenever

$$\sigma(1 - \sqrt{1 - 2\gamma}) < 0 \Leftrightarrow \gamma < 0. \quad (\text{A.20})$$

$R_1^{\text{GP}}(\gamma, \sigma, \theta)$ is constant in θ , whenever

$$\sigma(1 - \sqrt{1 - 2\gamma}) = 0 \Leftrightarrow \gamma = 0. \quad (\text{A.21})$$

In this case, it turns out that $R_1^{\text{GP}}(\gamma, \sigma, \theta) = \sigma$.

□

A.4. Proof of Proposition 5

From the expression (16), it follows directly that $\bar{R}_1^{\text{GP}}(\gamma_k, \sigma_k, \theta)$ increases with σ_k and with γ_k . Now, rewriting the risk measure as

$$\begin{aligned} \bar{R}_1^{\text{GP}}(\gamma_k, \sigma_k, \theta) = & \theta \frac{\frac{1-\gamma_k}{\sigma_k} \sum_{\substack{i=1 \\ i \neq k}}^{I+1} \frac{\sigma_i(\sqrt{1-2\gamma_i} - \sqrt{1-2\gamma_k})}{(1-\gamma_i)\sqrt{1-2\gamma_i}}}{\left(1 + \frac{(1-\gamma_k)\sqrt{1-2\gamma_k}}{\sigma_k} \sum_{\substack{i=1 \\ i \neq k}}^{I+1} \frac{\sigma_i}{(1-\gamma_i)\sqrt{1-2\gamma_i}} \right) \left(1 + \frac{1-\gamma_k}{\sigma_k} \sum_{\substack{i=1 \\ i \neq k}}^{I+1} \frac{\sigma_i}{1-\gamma_i} \right)} \\ & + \frac{1}{1 + \frac{1-\gamma_k}{\sigma_k} \sum_{\substack{i=1 \\ i \neq k}}^{I+1} \frac{\sigma_i}{1-\gamma_i}}, \end{aligned}$$

since $\frac{1-\gamma_k}{\sigma_k} > 0$ and the denominator of the factor multiplying θ is positive, one concludes that $\bar{R}_1^{\text{GP}}(\gamma_k, \sigma_k, \theta)$ increases (decreases) with θ if the quantity in the numerator denoted by

$$F(\gamma_k) = \sum_{\substack{i=1 \\ i \neq k}}^{I+1} \frac{\sigma_i(\sqrt{1-2\gamma_i} - \sqrt{1-2\gamma_k})}{(1-\gamma_i)\sqrt{1-2\gamma_i}}$$

is positive (negative); $\bar{R}_1^{\text{GP}}(\gamma_k, \sigma_k, \theta)$ remains constant if this quantity vanishes.

□

A.5. Proof of Proposition 6

From the definition it follows directly that $R_2^{\text{GP}}(\gamma, \sigma, \theta; p)$ increases with σ for $\gamma \neq 0$ and for $\gamma = 0$.

Now, consider the case $\gamma \neq 0$. Calculating the derivative of $R_2^{\text{GP}}(\gamma, \sigma, \theta; p)$ with respect to γ , one obtains

$$\frac{dR_2^{\text{GP}}(\gamma, \sigma, \theta; p)}{d\gamma} = \sigma \left[\frac{\theta(1 - p^{-\gamma} - \gamma p^{-\gamma} \log(p))}{\gamma^2} + \frac{(1 - \theta)}{(1 - \gamma)^2} \right]. \quad (\text{A.22})$$

Analysing the expression in the numerator of the first term of the sum in (A.22), given by

$$g(\gamma; p) = 1 - p^{-\gamma} - \gamma p^{-\gamma} \log(p) = 1 - p^{-\gamma} [1 + \gamma \log(p)], \quad (\text{A.23})$$

for $p \in (0, 0.05]$, one deduces the following results.

1. Let $0 < \gamma < \frac{1}{2}$. Proving that $g(\gamma; p) > 0$ is equivalent to prove that $e^x(1 - x) < 1$, where $x = -\gamma \log(p) > 0$, for $p \in (0, 0.05]$. Using Taylor expansion of e^x , one has

$$e^x(1 - x) = \left(1 + x + \frac{x^2}{2} + \frac{x^3}{3!} + \dots \right) (1 - x) = 1 - \frac{x^2}{2} - \dots,$$

where all remaining terms are subtracted and positive. Thus, $e^x(1 - x) < 1$, so that one can conclude that $g(\gamma; p) > 0$.

2. Let $\gamma < 0$, then $g'(\gamma; p) = \gamma p^{-\gamma} (\log p)^2 < 0$, so that $g(\gamma; p)$ is strictly decreasing. Furthermore, since $\lim_{\gamma \rightarrow -\infty} g(\gamma; p) = \lim_{\gamma \rightarrow -\infty} \left(1 - \frac{1 + \log p^\gamma}{p^\gamma}\right) = 1$ and $g(0; p) = 0$ for $p \in (0, 0.05]$, it follows that $g(\gamma; p) > 0$.

Therefore, and from the fact that the term on the right hand side of the sum in (A.22) is positive, one concludes that $\frac{dR_2^{\text{GP}}(\gamma, \sigma, \theta; p)}{d\gamma} > 0$. Hence, $R_2^{\text{GP}}(\gamma, \sigma, \theta; p)$ increases with γ .

Rewriting $R_2^{\text{GP}}(\gamma, \sigma, \theta; p)$ for $\gamma \neq 0$ as

$$R_2^{\text{GP}}(\gamma, \sigma, \theta; p) = \theta \sigma \left[\frac{p^{-\gamma} - 1}{\gamma} - \frac{1}{1 - \gamma} \right] + \frac{\sigma}{1 - \gamma} = \theta \sigma f(\gamma; p) + \frac{\sigma}{1 - \gamma}, \quad (\text{A.24})$$

where

$$f(\gamma; p) = \frac{p^{-\gamma} - 1}{\gamma} - \frac{1}{1 - \gamma}, \quad (\text{A.25})$$

and analysing $f(\gamma; p)$ for $p \in (0, 0.05]$, one obtains the following results.

1. Let $0 < \gamma < \frac{1}{2}$. Differentiating (A.25) with respect to p one gets $f_p(\gamma; p) = -p^{-\gamma-1}$, which is negative. Therefore, $f(\gamma; p)$ is strictly decreasing in p and one has $f(\gamma; p) \geq f(\gamma; 0.05)$ for $p \in (0, 0.05]$. In order to show that $f(\gamma; p) > 0$ it is sufficient to prove that $f(\gamma; 0.05) > 0$. From (A.25) one obtains

$$f(\gamma; 0.05) = \frac{20^\gamma - 1}{\gamma} - \frac{1}{1 - \gamma}.$$

Proving that $f(\gamma; 0.05) > 0$ is equivalent to prove that

$$g(\gamma) = (20^\gamma - 1)(1 - \gamma) - \gamma > 0.$$

Now, $g(0) = 0$ and $g'(\gamma) = 20^\gamma [(1 - \gamma) \log(20) - 1]$. Since

$$(1 - \gamma) \log(20) > \frac{1}{2} \log(20) > 1,$$

one concludes that $g'(\gamma) > 0$ and, thus, $g(\gamma)$ is increasing for $0 < \gamma < \frac{1}{2}$. Due to the fact that $g(0) = 0$ and $g'(\gamma) > 0$, one has $g(\gamma) > 0$ for $0 < \gamma < \frac{1}{2}$. This permits concluding that $f(\gamma; 0.05) > 0$.

2. Let $\gamma < 0$. Differentiating $f(\gamma; p)$ with respect to p gives $f_p(\gamma; p) = -p^{-\gamma-1} < 0$, so that $f(\gamma; p)$ is strictly decreasing in p . Consequently, for $p \in (0, 0.05]$, $f(\gamma; p) \geq f(\gamma; 0.05)$. Thus, it is sufficient to prove that $f(\gamma; 0.05) > 0$. Now, let $\delta = -\gamma$. Then, $p^{-\gamma} = p^\delta$, and rewriting (A.25) accordingly, we define

$$F(\delta; p) = \frac{1 - p^\delta}{\delta} - \frac{1}{1 + \delta},$$

with $\delta > 0$. We will show that $F(\delta; p) > 0$, where it is sufficient to prove that $F(\delta; 0.05) > 0$. Note that for $p \in (0, 0.05]$, p^δ decreases as δ increases. Proving that $F(\delta; p) > 0$ is equivalent to prove

$$\frac{1 - p^\delta}{\delta} > \frac{1}{1 + \delta} \Leftrightarrow (1 + \delta)p^\delta < 1.$$

However, since $p^\delta \leq 0.05^\delta$, the last inequality holds if $(1 + \delta)0.05^\delta < 1$. Consider $g(\delta) = (1 + \delta)0.05^\delta$ and define $h(\delta) = \log(g(\delta)) = \log(1 + \delta) + \delta \log(0.05)$. Since $h'(\delta) = \frac{1}{1 + \delta} + \log(0.05) < 0$, for all $\delta > 0$, because $\log(0.05) \approx -2.9957$, and $\frac{1}{1 + \delta} < 1$, and due to the fact that $h(0) = 0$, then $h(\delta) = \log(g(\delta)) < 0$ for all $\delta > 0$. Hence, $g(\delta) < 1$ and, thus, $(1 + \delta)p^\delta < 1$, implying $F(\delta; p) > 0$. With $\gamma = -\delta$, it turns out that $f(\gamma; p) > 0$ for all $\gamma < 0$.

One concludes that $\frac{dR_2^{\text{GP}}(\gamma, \sigma, \theta; p)}{d\theta} > 0$, so that $R_2^{\text{GP}}(\gamma, \sigma, \theta; p)$ increases with θ for $0 < \gamma < \frac{1}{2}$ and for $\gamma < 0$.

Now, consider $\gamma = 0$. Rewriting $R_2^{\text{GP}}(\gamma, \sigma, \theta; p)$ as

$$R_2^{\text{GP}}(\gamma, \sigma, \theta; p) = \theta \sigma (-\log p - 1) + \sigma, \quad (\text{A.26})$$

since for $p \in (0, 0.05]$ one has $-\log(p) > 1$ (note that $\log(0.05) \approx -2.9957$), one concludes that $R_2^{\text{GP}}(\gamma, \sigma, \theta; p)$ increases with θ for $\gamma = 0$ and for $p \in (0, 0.05]$. □

A.6. Proof of Proposition 7

Consider the case $\gamma_k \neq 0$. From the definition it follows directly that $\overline{R}_2^{\text{GP}}(\gamma_k, \sigma_k, \theta; p)$ increases with σ_k . Calculating the derivative of $\overline{R}_2^{\text{GP}}(\gamma_k, \sigma_k, \theta; p)$ with respect to γ_k , one gets

$$\frac{d\overline{R}_2^{\text{GP}}(\gamma_k, \sigma_k, \theta; p)}{d\gamma_k} = \theta \frac{\frac{1 - p^{-\gamma_k} - \gamma_k p^{-\gamma_k} \log(p)}{\sigma_k (p^{-\gamma_k} - 1)^2} \Pi}{\left[1 + \frac{\gamma_k}{\sigma_k (p^{-\gamma_k} - 1)} \Pi\right]^2} + (1 - \theta) \frac{\frac{1}{\sigma_k} \sum_{i=1}^{I+1} \frac{\sigma_i}{1 - \gamma_i}}{\left[1 + \frac{1 - \gamma_k}{\sigma_k} \sum_{i=1}^{I+1} \frac{\sigma_i}{1 - \gamma_i}\right]^2}.$$

By inspection and from the analysis of the expression in the numerator in the first term of the sum, $1 - p^{-\gamma_k} - \gamma_k p^{-\gamma_k} \log(p)$, for $p \in (0, 0.05]$ presented in the proof of [Proposition 6](#), one concludes that $\frac{d\bar{R}_2^{\text{GP}}(\gamma_k, \sigma_k, \theta; p)}{d\gamma_k} > 0$, and, therefore, $\bar{R}_2^{\text{GP}}(\gamma_k, \sigma_k, \theta; p)$ increases with γ_k .

Rewriting $\bar{R}_2^{\text{GP}}(\gamma_k, \sigma_k, \theta; p)$ as

$$\bar{R}_2^{\text{GP}}(\gamma_k, \sigma_k, \theta; p) = \frac{\theta}{\sigma_k} \frac{(1 - \gamma_k) \sum_{i=1}^{I+1} \frac{\sigma_i}{1 - \gamma_i} - \frac{\gamma_k}{(p^{-\gamma_k} - 1)} \Pi}{\left[1 + \frac{\gamma_k}{\sigma_k (p^{-\gamma_k} - 1)} \Pi \right] \left[1 + \frac{1 - \gamma_k}{\sigma_k} \sum_{i=1}^{I+1} \frac{\sigma_i}{1 - \gamma_i} \right]} + \frac{1}{1 + \frac{1 - \gamma_k}{\sigma_k} \sum_{i=1}^{I+1} \frac{\sigma_i}{1 - \gamma_i}},$$

it turns out that $\bar{R}_2^{\text{GP}}(\gamma_k, \sigma_k, \theta; p)$ increases (decreases) with θ if the numerator of the factor multiplying θ is positive (negative) and is constant if it vanishes, leading to the stated result.

Consider the case $\gamma_k = 0$. From the definition one can observe that $\bar{R}_2^{\text{GP}}(\gamma_k, \sigma_k, \theta; p)$ increases with σ_k . The risk measure can be rewritten as

$$\bar{R}_2^{\text{GP}}(\sigma_k, \theta; p) = \frac{\theta}{\sigma_k} \frac{\sum_{i=1}^{I+1} \frac{\sigma_i}{1 - \gamma_i} + \frac{1}{\log p} \Pi}{\left[1 + \frac{1}{-\sigma_k \log p} \Pi \right] \left[1 + \frac{1}{\sigma_k} \sum_{i=1}^{I+1} \frac{\sigma_i}{1 - \gamma_i} \right]} + \frac{1}{1 + \frac{1}{\sigma_k} \sum_{i=1}^{I+1} \frac{\sigma_i}{1 - \gamma_i}},$$

consequently, $\bar{R}_2^{\text{GP}}(\sigma_k, \theta; p)$ increases (decreases) with θ if the numerator of the factor multiplying θ is positive (negative) and is constant if it vanishes leading to the established result. □

Appendix B. Supplementary tables and figures

Table B.7
Normalized GP risk measure components.

Month	$\overline{\text{SD}}[X]$	$\overline{\mathbb{E}}[X]$	$\overline{\text{VaR}}[X; 0.01]$	$\overline{\text{VaR}}[X; 0.05]$
Jan	0.1325	0.1431	0.1294	0.1379
Feb	0.1308	0.1326	0.1310	0.1325
Mar	0.1166	0.1153	0.1177	0.1167
Apr	0.1050	0.0983	0.1076	0.1021
May	0.0517	0.0513	0.0521	0.0518
Jun	0.0475	0.0396	0.0492	0.0429
Jul	0.0306	0.0303	0.0309	0.0307
Aug	0.0473	0.0407	0.0489	0.0437
Sep	0.0535	0.0426	0.0553	0.0468
Oct	0.0777	0.0769	0.0784	0.0778
Nov	0.0816	0.0927	0.0778	0.0865
Dec	0.0947	0.1063	0.0907	0.1000

Table B.8
Factor $F(\gamma_k)$ of [Proposition 5](#) for normalized Mean-SD risks.

Month	$F(\gamma_k)$	Month	$F(\gamma_k)$
Jan	-7.4075	Jul	1.3895
Feb	-2.0853	Aug	12.3212
Mar	0.7032	Sep	18.6351
Apr	5.4426	Oct	2.3481
May	0.9338	Nov	-11.4779
Jun	14.6283	Dec	-11.5259

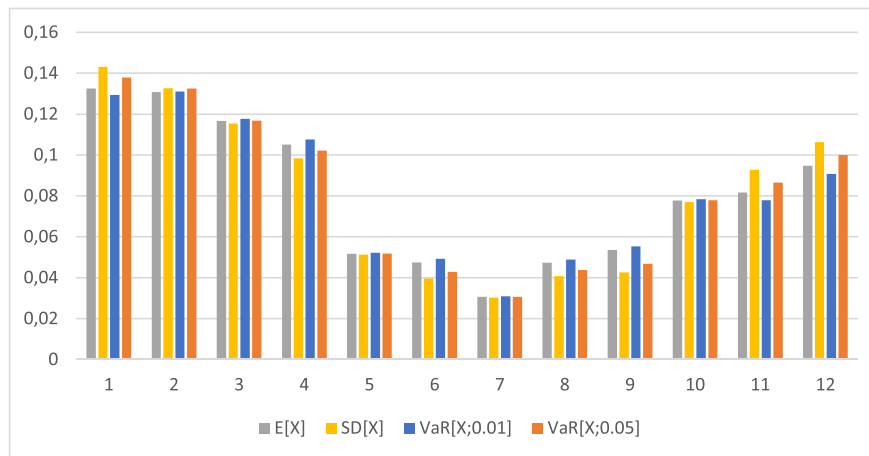


Fig. B.4. Mean, standard deviation and value at risk for $p = 0.01$ and $p = 0.05$ of $PM_{2.5}$ from GP modeling in each month.

Table B.9

Factor $G(\gamma_k)$ of Proposition 7 for normalized Mean-VaR risks for $p = 0.01$.

Month	$G(\gamma_k)$	Month	$G(\gamma_k)$
Jan	-10.74965	Jul	1.8505
Feb	-1.1313	Aug	13.7137
Mar	1.8712	Sep	17.5734
Apr	7.6213	Oct	1.8503
May	1.47888	Nov	-20.6830
Jun	15.4145	Dec	-18.2902

Table B.10

Factor $G(\gamma_k)$ of Proposition 7 for normalized Mean-VaR risks for $p = 0.05$.

Month	$G(\gamma_k)$	Month	$G(\gamma_k)$
Jan	-3.8028	Jul	1.0889
Feb	-0.0575	Aug	5.5244
Mar	1.0968	Sep	6.8760
Apr	3.2788	Oct	1.0888
May	0.9464	Nov	-7.7282
Jun	6.1301	Dec	-6.7780

Appendix C. Normalized GP Mean-VaR risk model for $p = 0.05$

The GP Mean-VaR risk model with normalized risk components $\overline{VaR}[X;0.05]$ and $\overline{E}[X]$, listed in Table B.7 (see Appendix B), is applied to the risk assessment of the twelve months. The risk models are presented graphically in Fig. C.5 for $\theta \in [0, 1]$.

The values of the factors $G(\gamma_k)$ of Proposition 7, that are displayed in Table B.10 (see Appendix B), permit again concluding that for January, February, November and December the risk decreases with θ since $G(\gamma_k) < 0$. For the remaining months the risk increases with θ , since $G(\gamma_k) > 0$. This behaviour of the risk lines can be confirmed in Fig. C.5.

From the graphical representations one can already observe that the lines exhibit, in general, a relatively constant trend, with only two lines appearing to intersect. The change in θ only affects the ordering of April and December, where for $\theta < 0.7921$ December is classified as fourth riskiest month and April as fifth riskiest month and for $\theta \in [0.7921, 1]$ these months change their risk positions.

One can conclude that also the normalized Mean-VaR risk model with $p = 0.05$ classifies the cold period from November to April as riskiest period for air pollution caused by $PM_{2.5}$, where January is classified as riskiest month for $\theta \in [0, 1]$ and February as second riskiest month. July has the lowest risk for $\theta \in [0, 1]$, and June is classified as presenting the second lowest risk.

Table C.11

Mean-VaR risk ordering ($p = 0.05$) of months for interval ranges $I_\theta \subset [0, 1]$.

I_θ	Risk ordering
$[0, 0.7921)$	Jan > Feb > Mar > Dec > Apr > Nov > Oct > May > Sep > Aug > Jun > Jul
$[0.7921, 1]$	Jan > Feb > Mar > Apr > Dec > Nov > Oct > May > Sep > Aug > Jun > Jul

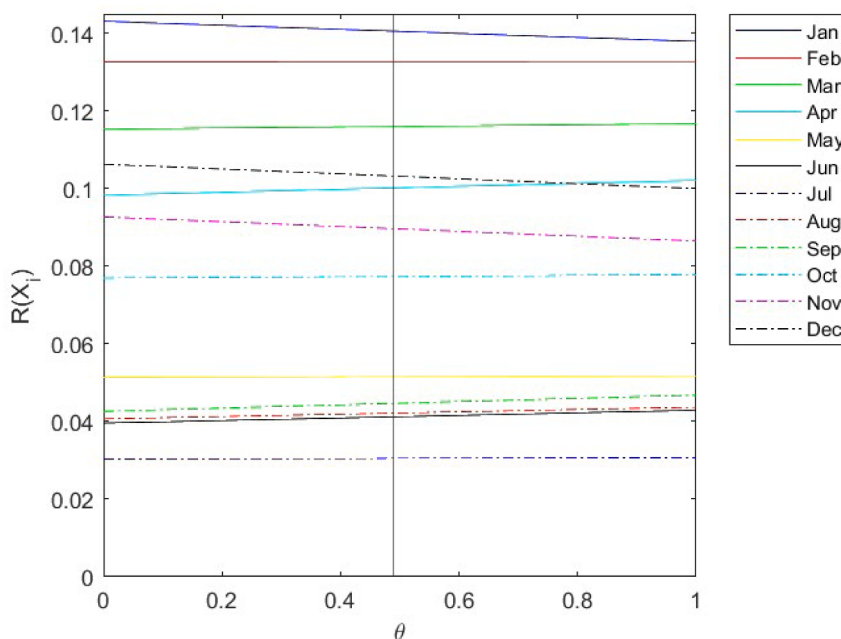


Fig. C.5. Normalized Mean-VaR risks for $p = 0.05$, where the vertical line corresponds to $\theta = 0.4887$ obtained in (C.1).

Table C.12

Optimality values S_i , utility degrees k_i , rank orders using ARAS with mean and VaR with $p = 0.05$ and standard deviation weighting procedure (the rank order 1 corresponds to the best classified, i.e. least risky, alternative and the rank order 12 to the worst classified alternative, i.e. the one with the highest risk) and Bootstrap resampling results ($SD(S_i)$, MRP, MR, $SD(R)$, MedR).

Month	S_i	k_i	rank	$SD(S_i)$	MRP	MR	$SD(R)$	MedR
Jan	0.1406	0.2169	12	0.0275	11.3895	12	0.7687	12
Feb	0.1325	0.2300	11	0.0269	10.9962	11	0.8885	11
Mar	0.1160	0.2628	10	0.0251	10.2227	10	0.9759	10
Apr	0.1001	0.3045	8	0.0207	8.6235	9	0.8857	9
May	0.0515	0.5915	5	0.0102	4.4616	5	0.7344	5
Jun	0.0413	0.7390	2	0.1693	3.4373	3	1.9014	3
Jul	0.0305	1.0000	1	0.0063	1.0200	1	0.1643	1
Aug	0.0422	0.7232	3	0.0089	2.7632	2	0.8558	3
Sep	0.0447	0.6825	4	0.0100	3.5535	4	1.0097	4
Oct	0.0773	0.3943	6	0.0156	6.1634	6	0.5132	6
Nov	0.0896	0.3401	7	0.0167	7.0486	7	0.6524	7
Dec	0.1032	0.2954	9	0.0201	8.3205	8	0.9116	8

ARAS with mean and VaR with $p = 0.05$ criteria

Considering the ARAS method with risk criteria mean and VaR with $p = 0.05$, where again the standard deviation weighting method is applied for computing the criteria weights, yielding

$$\tilde{\theta} = w_1 = 0.4887, \quad 1 - \tilde{\theta} = w_2 = 0.5113, \quad (\text{C.1})$$

one gets the results in Table C.12. As can be seen, the ARAS ranking aligns with the risk ordering in Table C.11 for the interval $[0, 0.7921]$, since $\theta = 0.4887$.

The sensitivity analysis based on variations in the weighting factor (see Table C.11), indicates that this risk classification method is highly robust, as only one rank reversal occurs.

Uncertainty analysis for the normalized Mean-VaR risk model with $p = 0.05$

Using the resampling-based framework for uncertainty quantification, previously introduced and applied to the other two risk models, we evaluated the normalized risk S_i , for $i = 1, \dots, 12$, as defined in (24), across 10,000 bootstrap samples, now under the Mean-VaR model with $p = 0.05$. Table C.12 reports the median ranking (MedR) obtained from the 10,000 Bootstrap rankings, and also the mean ranking positions (MRP) along with the corresponding mean ranking (MR). The median ranking MedR produces only one tie, between June and August. Solving this tie using the mean ranking yields a ranking identical to that obtained with MR. Thus, the final Bootstrap ranking is given by MR (derived from MedR with the tie broken using the mean ranking).

Similar to the Mean-SD risk model we observe the same interchanging ranking positions between April and December and between June and August. Spearman and Kendall analyses again yield coefficients in the interval $(0.93, 1]$ with p -values below 0.01, indicating

strong and statistically significant agreement among the rankings. Ranking uncertainty was quantified using bootstrap-based standard deviations of both the estimated scores S_i and the ranks, respectively denoted by $SD(S_i)$ and $SD(R)$. These are also reported in Table C.12. The variability of S_i is heterogeneous, with June exhibiting substantially higher variability ($SD(S_i) = 0.1693$) than all other months, indicating pronounced instability in its underlying score, while the remaining months display moderate to low variability, particularly July and May. The bootstrap variability of the rankings, measured by $SD(R)$, is generally low, suggesting a high degree of overall ranking stability. In particular, July shows very strong stability ($SD(R) = 0.1643$), followed by October and November, whereas June again stands out as the most uncertain month ($SD(R) = 1.9014$), reflecting sensitivity to resampling. Overall, these results confirm that ranking uncertainty is driven by both the variability and proximity of the underlying scores, with months characterized by more distinct S_i values, such as July, exhibiting highly robust rankings, while intermediate months remain more sensitive to sampling variability.

References

- Aksarayli, M., Pala, O., 2018. A polynomial goal programming model for portfolio optimization based on entropy and higher moments. *Expert Syst. Appl.* 94, 185–192.
- Aven, T., 2016. Risk assessment and risk management: review of recent advances on their foundation. *Eur. J. Oper. Res.* 253, 1–13.
- Balkema, A.A., Haan, L.D., 1974. Residual life time at great age. *Ann. Probab.* 2, 792–804.
- Beirlant, J., Goegebeur, Y., Segers, J., Teugels, J., 2004. *Statistics of Extremes: Theory and Applications*. Wiley, Chichester.
- Brachinger, H.W., Weber, M., 1997. Risk as a primitive: a survey of measures of perceived risk. *Oper. Res. Spektrum* 19, 235–250.
- Bruto, I., 2020. A decision model based on expected utility, entropy and variance. *Appl. Math. Comput.* 379, 125285.
- Bruto, I., 2022. The normalized expected utility-entropy and variance model for decisions under risk. *Int. J. Approx. Reason.* 148, 174–201.
- Bruto, I., 2025. A multi-risk criteria assessment methodology applied to the analysis of nitrogen dioxide concentrations in european capital cities. *Environ. Ecol. Stat.* 32, 827–854. <https://doi.org/10.1007/s10651-025-00666-6>
- Castillo, J.D., Padilla, M., 2016. Modelling extreme values by the residual coefficient of variation. *Stat. Oper. Res. Trans.* 40 (2), 303–320. <https://raco.cat/index.php/SORT/article/view/316240>.
- Castillo, J.D., Serra, I., 2015. Likelihood inference for generalized pareto distribution. *Comput. Stat. Data Anal.* 83, 116–128.
- Castillo, J.D., Soler, D.M., Serra, I., 2019. ercv: Fitting Tails by the Empirical Residual Coefficient of Variation. R Package Version 1.0.1. <https://CRAN.R-project.org/package=ercv>.
- Chatterjee, S., Chakraborty, S., 2024. A study on the effects of objective weighting methods on tosis-based parametric optimization of non-traditional machining processes. *Decis. Anal. J.* 11, 100451.
- Core, T., 2020. R: a language and environment for statistical computing. In: R Foundation for Statistical Computing. Vienna, Austria <https://www.R-project.org/>.
- Davison, A.C., Smith, R.L., 1990. Models for exceedances over high thresholds (with discussion). *J. R. Stat. Soc. B* 52, 393–442.
- European Environment Agency, 2025. Air quality data. <https://www.eea.europa.eu/en/datahub>.
- EEA, 2024. Europe's air quality status 2024. Briefing no. 06/2024. <https://doi.org/10.2800/5970>
- Ferreira, M., Gomes, M.I., Leiva, V., 2012. On an extreme value version of the birnbaum-saunders distribution. *REVSTAT-Stat. J.* 10 (2), 181–210. <https://doi.org/10.57805/revstat.v10i2.116>
- Geissel, S., Graf, H., Herbinger, J., 2022. Portfolio optimization with optimal expected utility risk measures. *Ann. Oper. Res.* 309, 59–77. <https://doi.org/10.1007/s10479-021-04403-7>
- Hwang, C.-L., Lai, J.-Y., Liu, T.-Y., 1993. A new approach for multiple objective decision making. *Comput. Oper. Res.* 20, 889–899. [https://doi.org/10.1016/0305-0548\(93\)90109-V](https://doi.org/10.1016/0305-0548(93)90109-V)
- Jia, J., Dyer, J.S., 2009. Decision making based on risk-value tradeoffs. In: Brams, S.J., Gehrlein, W.V., Roberts, F.S. (Eds.), *The Mathematics of Preference, Choice and Order. Studies in Choice and Welfare*. Springer.
- Kaas, R., Goovaerts, M., Dhaene, J., Denuit, M., 2008. *Modern Actuarial Risk Theory*. Heidelberg, Springer.
- Kahneman, D., Tversky, A., 1979. Prospect theory: an analysis of decisions under risk. *Econometrica* 47, 263–290.
- Klugman, S., Panjer, H., Willmot, G., 2019. *Loss Models: from Data to Decisions*. Wiley, New York.
- Leiva, V., Ferreira, M., Gomes, M.I., Lillo, C., 2016. Extreme value birnbaum-saunders regression models applied to environmental data. *Stoch. Environ. Res. Risk Assess.* 30, 1045–1058. <https://doi.org/10.1007/s00477-015-1069-6>
- Li, X., Qin, Z., Kar, K., 2010. Mean-variance-skewness model for portfolio selection with fuzzy returns. *Eur. J. Oper. Res.* 202 (1), 239–247. <https://doi.org/10.1016/j.jejor.2009.05.003>
- Lin, Y.-C., Shih, H.-S., Lai, C.-Y., 2022. Classification of air quality zones and fine particulate matter sensitive areas by risk assessment approach. *Environ. Res.* 215, 114208. <https://doi.org/10.1016/j.envres.2022.114208>
- Liu, J., Scaife, A.A., Dunstone, N., 2023. Predictability and risk of extreme winter PM2.5 concentration in Beijing. *J. Meteorol. Res.* 37, 632–642. <https://doi.org/10.1007/s13351-023-3023-8>
- Liu, N., Xu, Z., 2021. An overview of aras method: theory development, application extension, and future challenge. *Int. J. Intell. Syst.* 36, (1) 3524–3565. <https://doi.org/10.1002/int.22425>
- Lomba, J.S., Alves, M.I., 2020. L-moments for automatic threshold selection in extreme value analysis. *Stoch. Environ. Res. Risk Assess.* 34, 465–491. <https://doi.org/10.1007/s00477-020-01789-x>
- Markowitz, H., 2000. *Mean-Variance Analysis in Portfolio Choice and Capital Markets*. Wiley.
- Martins, L.D., Wikuats, C. F.H., Capucim, M.N., Almeida, D.S., Costa, S. C.D., Albuquerque, T. R. D.A., Carvalho, V. S.B., Freitas, E.D., Andrade, M.F., Martins, J.A., 2017. Extreme value analysis of air pollution data and their comparison between two large urban regions of south america. *Weath. Clim. Extrem.* 18, 44–54. <https://doi.org/10.1016/j.wace.2017.10.004>
- Moreira, E., Ferreira, M., 2025. Finding the tail of a distribution: analysis of a method based on the coefficient of variation. *J. Appl. Stat.* pp. 1–24. <https://doi.org/10.1080/02664763.2025.2487540>
- von Neumann, J., Morgenstern, O., 1944. *Theory of Games and Economic Behavior*. Princeton, Princeton University Press.
- Orellano, P., Reynoso, J., Quaranta, N., 2020. Short-term exposure to particulate matter (PM10 and PM2.5), nitrogen dioxide (NO2), and ozone (O3) and all-cause and cause-specific mortality: systematic review and meta-analysis. *Environ. Int.* 142, 105876. <https://doi.org/10.1016/j.envint.2020.105876>
- Pickands, J. III, 1975. Statistical inference using extreme order statistics. *Ann. Stat.* 3, 119–131.
- Pollatsek, A., Tversky, A., 1970. A theory of risk. *J. Math. Psychol.* 7, 540–553.
- Sahoo, S., Goswami, S., 2023. A comprehensive review of multiple criteria decision-making (mcgm) methods: advancements, applications, and future directions. *Decis. Mak. Adv.* 1, 25–48. <https://doi.org/10.31181/dma1120237>
- Usta, I., Kantar, Y.M., 2011. Mean-variance-skewness-entropy measures: a multi-objective approach for portfolio selection. *Entropy* 13, 117–133. <https://doi.org/10.3390/e1301011>
- Vavrek, R., 2019. Evaluation of the impact of selected weighting methods on the results of the tosis technique. *Int. J. Inf. Technol. Decis. Mak.* 18, 1821–1843. <https://doi.org/10.1142/S021962201950041X>
- Zavadskas, E.K., Turskis, Z., 2010. A new additive ratio assessment (ARAS) method in multicriteria decision-making, *Ukio Technologinis ir Ekonominis Vystymas*. 16, 159–172 <https://doi.org/10.3846/tede.2010.10>