

Relatively Robust Multicriteria Decisions

Thomas A. Weber^a

^a Chair of Operations, Economics and Strategy, École Polytechnique Fédérale de Lausanne, CH-1015 Lausanne, Switzerland

Contact: thomas.weber@epfl.ch,  <https://orcid.org/0000-0003-2857-4754> (TAW)

Received: February 7, 2025

Revised: May 26, 2025

Accepted: May 28, 2025

Published Online in Articles in Advance:
August 14, 2025

<https://doi.org/10.1287/mnsc.2025.00510>

Copyright: © 2025 The Author(s)

Abstract. For a general multicriteria decision problem with linear scalarization and unknown weights, we propose relatively robust decisions, which are Pareto-efficient and at the same time maximize a performance index. The latter measures the worst-case ratio, attained by the weighted objective relative to its maximum value, with respect to all possible weights. The main results include a simple boundary representation of the performance index as the minimum of criterion-specific performance ratios, and a computationally simple method of determining a relatively robust decision up to any prespecified performance tolerance by maximizing an ε -augmented performance index. The proposed method relies merely on the continuity of all criterion functions and the compactness of the set of feasible decisions which may be nonconvex. This imposes no restrictions at all for any finite action set. A notable feature of our method is that it endogenously yields the tradeoffs between all criteria, including a performance guarantee relative to decisions justified by any other weighting. A number of structural results, examples, and applications are provided, as well as generalizations to allow for limited weight ambiguity, criterion ambiguity, and generalized aggregation of criteria based on an axiomatic foundation.

History: Accepted by Peng Sun, optimization and decision analytics.



Open Access Statement: This work is licensed under a Creative Commons Attribution-NoDerivatives 4.0 International License. You are free to download this work and share with others commercially or noncommercially, but cannot change in any way, and you must attribute this work as “*Management Science*. Copyright © 2025 The Author(s). <https://doi.org/10.1287/mnsc.2025.00510>, used under a Creative Commons Attribution License: <https://creativecommons.org/licenses/by-nd/4.0/>.”

Keywords: criterion ambiguity • multicriteria decisions • relative regret • robust optimization • weight ambiguity

1. Introduction

In real-world decision making, evaluating alternatives often involves multiple, sometimes conflicting criteria. Whether considering investment portfolios under Environmental, Social, and Governance (ESG) parameters, selecting products based on bundles of attributes, or valuing companies for both profitability and sustainability, decision makers must navigate tradeoffs between competing objectives. Similarly, lifecycle environmental impact assessments—such as comparing vehicles with electric, combustion, or hybrid engines—require reconciling diverse metrics like emissions, cost, and resource consumption. These complex but often inevitable comparisons highlight the critical need for robust multicriteria optimization frameworks.

Multicriteria optimization involves identifying solutions that balance conflicting objectives in a manner consistent with the decision maker’s priorities. Traditional approaches often rely on scalarization techniques, where multiple criteria are combined into a single weighted objective function using weighted sums. However, these methods depend heavily on the precise specification of weights, which are rarely known a priori and can be difficult to justify. This uncertainty

complicates the search for decisions that are both Pareto-efficient and robust to variations in tradeoff preferences. To address these challenges, we specialize the concept of relatively robust decisions by Weber (2023) so as to refer to decisions that achieve Pareto-efficiency among all relevant criteria while also maximizing a performance index designed to account for the ambiguity in weights. Specifically, the performance index measures the worst-case (WC) ratio attained by the weighted objective relative to its maximum value over all possible weight configurations. By focusing on worst-case performance, this framework provides guarantees of robustness and tradeoff transparency, which are particularly valuable in high-stakes or uncertain decision contexts.

1.1. Practical Examples

The approach developed here can be used for virtually all multicriteria decision problems, such as the following three example applications.¹

- **ESG Investing:** Investors face the challenge of balancing financial returns with social and environmental impact. For example, a fund manager might evaluate portfolios based on criteria such as profitability, carbon

footprint, and diversity inclusion. Weighting these criteria is inherently subjective, and the optimal portfolio might vary widely depending on the chosen weights. The relatively robust optimization framework proposed here enables the identification of investment strategies that remain robust across different weight configurations, offering a performance guarantee regardless of the specific preferences.²

- **Lifecycle Assessment of Vehicles:** Consider evaluating the environmental impact of electric, combustion, and hybrid cars. Criteria might include greenhouse gas emissions, energy efficiency, and resource use (see, e.g., Hawkins et al. 2012). A relatively robust decision could pinpoint vehicle types or designs that perform well across a broad range of plausible weightings, resulting in a balanced and defensible choice.

- **Product Evaluation and Design:** Companies frequently design products to optimize attributes such as cost, durability, and aesthetic appeal. For example, in designing a smartphone, decision makers must weigh the importance of battery life, screen quality, and price.³ Relatively robust multicriteria optimization helps determine design specifications that ensure competitive performance across various market segments with diverse preferences, valuing the availability of different attributes with different weights.

For the application of our method, one only needs that there exists a default action, that is, a baseline alternative that performs adequately across all criteria (which can always be achieved by reindexing evaluation scales), together with the technical assumption that criteria are continuous in actions, and that the finite-dimensional action set is closed and bounded (i.e., compact).

1.2. Literature

1.2.1. Origins of Multicriteria Decision Making. The idea of multicriteria optimization in Economics can be traced back to the distribution of resources among different individuals, leading to a set of undominated solutions such as the “contract curve” proposed by Edgeworth (1881, p. 21) for a simple exchange economy, and more generally a set of efficient outcomes as implied by Pareto (1894; 1897, sections 721–723),⁴ which cannot be improved upon for one agent without making another agent worse off. The latter avoids a direct interpersonal comparison of individuals’ utilities (or “ophelimities” in Pareto’s terminology); see also Harsanyi (1955) as well as Keeney and Raiffa (1993, chapter 10) who explore group utility functions. The drawback of such an agnostic approach to optimality with multiple criteria is that the set of Pareto-optimal allocations, because of its typically large cardinality, offers only imperfect guidance about which solution should actually be implemented. For example, in the two-agent exchange economy, the contract curve usually includes allocations that attribute all resources to any single

individual, which almost completely undermines the notion of *multicriteria* optimization.

In Operations Research, “goal programming” refers to the notion of minimizing the weighted deviation from criterion-specific targets (Charnes and Cooper 1961). The technique was first employed in the context of executive compensation based on different “factors” (i.e., criteria) using linear programming techniques (Charnes et al. 1955). This basic approach is usually applied to a “utopian” (or “ideal”) target point, which corresponds to the (generically infeasible) vector of individually maximized criteria. Depending on the distance measure (e.g., a weighted Chebyshev distance; see, e.g., Steuer 1986, chapter 14),⁵ the corresponding solutions trade off among criteria according to the specified weights, and they are naturally Pareto-efficient.⁶ Instead of minimizing the distance to the ideal point (in the outcome space), it is also possible to maximize the distance to a (generically infeasible) “nadir” point, which contains the minimum value of each individual criterion on the Pareto-efficient set.⁷ In this approach, known as “compromise programming” (Zeleny 1974), the appropriate choice of the weights for the different criteria remains the critical point, and in our view, very little satisfying progress has been made in the assignment of weights without imposing subjectivity, which arguably amounts to picking a solution from the Pareto-efficient set. For instance, based on an exogenous ranking of the criterion importance, it is possible to apply an ordered weighted average (Yager 1988) which in turn can be related to compromise programming (Zarghami and Szidarovszky 2010, Wang and Fu 2020).⁸ Besides being subjective from the start by requiring the imposition of a preference order on criteria, this method does not provide any nontrivial performance guarantee relative to other choices of weights and/or importance rankings that might be plausible for other decision makers.⁹

1.2.2. Relative Robustness. By contrast, we approach the selection of weights from the standpoint of relative robustness (Weber 2023), involving only comparisons between feasible points, thus avoiding fictitious targets such as the aforementioned utopian and nadir points. Rather than minimizing a fixed distance metric with specific weights, the proposed method evaluates performance in terms of the worst-case tradeoffs among criteria, ensuring a solution that remains defensible regardless of the exact weighting chosen ultimately. The underlying robustness measure, equivalent to relative regret, has been used in computer science to evaluate the relative performance of algorithms (Sleator and Tarjan 1985, Ben-David and Borodin 1994), for the scenario-based evaluation of operational decisions (Kouvelis and Yu 1997), in price discrimination (Han and Weber 2023), robust optimization (Weber 2024), and fair resource allocation (Goel et al. 2009). The idea of absolute regret (AR)

is due to Savage (1951), based on the maximin robustness approach by Wald (1945) in his general treatment of sequential decision problems. The main issue with absolute regret is that it is inherently sensitive to the scale of the reference point. This sensitivity may make it impossible to derive reasonable performance guarantees, such as ensuring positive profits in a monopoly pricing problem with unknown demand (Weber 2025), because a zero-profit reference point is intrinsically small-scale. A *relative* robustness approach, we argue, yields more acceptable results, particularly in the context of multicriteria optimization, where changing the units of any given criterion would usually affect solutions that are based on absolute performance measures.

1.2.3. Connection to Distributionally Robust Optimization. Because one can reinterpret a (normalized) weight vector as a probability distribution, our approach is naturally related to distributionally robust optimization (DRO), where the true probability distribution governing uncertain parameters is unknown but assumed to lie within a known ambiguity set. Important early contributions include Delage and Ye (2010), who studied DRO with moment-based sets, and Ben-Tal et al. (2013), who developed tractable reformulations for DRO with Wasserstein ambiguity sets. More recently, Blanchet and Murthy (2019) and Mohajerin Esfahani and Kuhn (2018) developed general formulations based on Wasserstein balls, offering strong out-of-sample guarantees. Whereas DRO typically focuses on expectations or risk measures over stochastic uncertainty, our approach generalizes worst-case robustness to a multicriteria setting without requiring a probabilistic model. This positions our proposed framework as a deterministic analogue to DRO, where ambiguity arises from unknown relative preferences and state-dependent criteria rather than unknown probability distributions.

1.2.4. Relative Evaluations. The proposed approach to robust multicriteria decisions is entirely relative, in the sense that the question underlying all of our developments is “How well am I doing relative to how well I could be doing?” Indeed, the idea that the size of an object can be judged only relative to other objects goes back at least to the Taoist writings of Zhuang Zhou in the fourth century BC.¹⁰ In Economics, Cournot (1838) was among the first to note that there is no absolute value (“Il n’y a pas de valeurs absolues,” p. 22) and that inference from a social system can be likened to the observation of astronomical objects and their relative positions to each other (pp. 15–16), concluding by analogy that the concept of value is fundamentally relative (“Il y en a en ce sens que des valeurs relatives,” p. 18). In fact, human perception is inherently relative, as demonstrated by Weber (1846) and his student Fechner (1860) in extensive experiments which showed that across

different senses (e.g., hearing, touch, and vision) the minimum perceptible difference is proportional to the current level of the stimulus, giving rise to the Weber-Fechner law of psychophysics. Similarly, in an economic context, it is often relative reference points such as one’s current wealth level (Kahneman and Tversky 1979) or the outcomes experienced by others (Loewenstein et al. 1989) that tend to determine human behavior. For example, when pondering whether to purchase a product from a cheaper store within walking distance, humans base decisions less on absolute gains than on the prospective relative savings in expenditure (Kahneman and Tversky 1984).

Beyond the aforementioned congruence with human perception, there are other practical arguments for relative evaluations. First, there is the *insensitivity to scale*, common to all relative measures such as internal rate of return (IRR), demand elasticity, profit margin, or relative regret, which allows for a direct comparison across different sizes. For example, with IRR one can readily benchmark projects of different financial magnitudes against outside investment options (of comparable risk), whereas the equivalent absolute indicator of net present value (NPV) remains silent about the required absolute investment (see, e.g., Weber 2014).¹¹ Second, normalization facilitates *fairness and equity*. When stakeholders have differing capacities or baselines, a relative comparison may help to ensure fairness (Goel et al. 2009). Third, relative evaluation criteria allow for *comparability across contexts*. For example, demand elasticity (as introduced by Marshall 1890, p. 162) is a unit-free relative measure that allows one to compare the changes of demand relative to price changes across widely differing goods, irrespective of the underlying unity of measurement (e.g., units of cars versus units of electric power). This last point is especially salient for multicriteria decision making, as the units (and inherent scale) for different criteria generally differ, so that a robustness criterion (and robust decision) should remain unaffected if the values of a given criterion are all multiplied by 10, for example. The relatively robust framework developed here uses a relative worst-case performance perspective. This methodology not only identifies solutions with reliable trade-off characteristics, but also provides a transparent representation of the tradeoffs themselves. These tradeoffs are reflected in a robust weight vector consistent with a robust solution.

1.3. Outline

The remainder of this paper is organized as follows: Section 2 discusses the multicriteria optimization problem and associated comparative statics, together with the performance index for a robust evaluation of different decisions. Section 3 introduces pseudo-robustness and Pareto-efficiency, which together characterize relatively robust decisions. Here we also provide a computational

approach for determining a relatively robust decision up to any given performance tolerance, and we allow for the possibility of close-to-arbitrary restrictions in the set of admissible weights, for example, based on a priori knowledge about physical constraints or a given priority ranking of criteria. In addition, there are extensions to criterion ambiguity and general aggregation of criteria, followed by a practical guide for how to apply the method. Section 4 focuses on discrete applications where our framework relies on virtually no assumptions, so the approach can be entirely data-driven. Section 5 concludes.

2. Basic Framework

Let $\mathcal{X} \subset \mathbb{R}^m$ be a nonempty, compact action set, for a given integer $m \geq 1$. Consider a decision maker who faces the multicriteria optimization problem of having to select an action (or decision, or point) $x \in \mathcal{X}$ so as to “simultaneously maximize” the continuous functions $f_i: \mathcal{X} \rightarrow \mathbb{R}_+$, for $i \in \mathcal{N} = \{1, \dots, n\}$, each of which is referred to as a *criterion*, where $n \geq 1$ is a given integer.

Remark 1. (i) The continuity of each criterion f_i ensures that small perturbations in the action x yield correspondingly small changes in the objective value. Notably, this assumption is trivially satisfied at any isolated point of $\mathcal{X}$.¹² In particular, this means that there is no imposed regularity requirement when the action set is finite. (ii) The requirement that each f_i be nonnegative is without loss of generality: any real-valued (continuous) criterion $\hat{f}_i$ can be translated as $f_i = \hat{f}_i - \hat{f}_i^\bullet \geq 0$, where $\hat{f}_i^\bullet = \min_{x \in \mathcal{X}} \hat{f}_i(x)$ denotes the minimum value of $\hat{f}_i$.¹³

2.1. Decision Problem

To evaluate goal achievement for any decision $x \in \mathcal{X}$, the decision maker considers a scalarization of his multicriteria optimization problem by means of a *weighted objective*,

$$F(x|\lambda) = \sum_{i=1}^n \lambda_i f_i(x), \quad x \in \mathcal{X}, \quad (1)$$

where $\lambda = (\lambda_1, \dots, \lambda_n) \in \Delta = \{\lambda \in \mathbb{R}_+^n : \|\lambda\|_1 = 1\}$ is a (normalized) vector of weights (with $\lambda_i \geq 0$ for all $i \in \mathcal{N}$, and $\lambda_1 + \dots + \lambda_n = 1$). Maximizing the weighted objective yields the set of ex post optimal decisions,

$$\mathcal{X}(\lambda) = \arg \max_{x \in \mathcal{X}} F(x|\lambda), \quad \lambda \in \Delta, \quad (2)$$

which by the extreme-value theorem (Rudin 1976, theorem 4.16, p. 89) is nonempty. Moreover, by the maximum theorem (Berge 1963, p. 116) the (set-valued) mapping $\mathcal{X}: \Delta \rightrightarrows \mathcal{X}$ is compact-valued and upper semicontinuous.

Remark 2. The weighted objective is homogeneous of degree one, that is, for all $\alpha > 0$ and $\lambda \in \Delta$, it is

$F(\cdot|\alpha\lambda) = \alpha F(\cdot|\lambda)$, whereas the set of ex post optimal decisions is homogeneous of degree zero, in the sense that $\mathcal{X}(\alpha\lambda) = \mathcal{X}(\lambda)$, for all $\alpha > 0$ and $\lambda \in \Delta$. It is therefore possible to extend the definition of $F(x|\cdot)$ in Equation (1) and the definition of $\mathcal{X}(\cdot)$ in Equation (2) to a domain containing any nonzero weight vector $w \in \mathbb{R}_+^n \setminus \{0\}$, because a unique normalized weight $\lambda = \alpha w \in \Delta$, with $\alpha = 1/(\sum_{i=1}^n w_i) > 0$, is always available.¹⁴

$$F(x|w) \triangleq (1/\alpha)F(x|\lambda), \quad x \in \mathcal{X}, \quad (1')$$

and

$$\mathcal{X}(w) \triangleq \mathcal{X}(\lambda) = \arg \max_{x \in \mathcal{X}} (1/\alpha)F(x|\lambda). \quad (2')$$

Remark 3. Limited ambiguity, that is, allowing for weights in a (nonempty, compact) subset of Δ , is discussed in Section 3.7. Criterion ambiguity is treated in Section 3.8, and Section 3.9 investigates the use of general multicriteria objectives. All three generalizations can be treated, after suitable adjustment, within the basic framework.

To keep matters nontrivial, we assume that there exists a (feasible) default decision (x_d) such that the decision maker’s weighted objective is positive, that is,

$$\exists x_d \in \mathcal{X}: F(x_d|\lambda) > 0, \quad \lambda \in \Delta. \quad (N)$$

The nontriviality condition (N) ensures that the ex post optimal objective (or value function) is positive:

$$F^*(\lambda) = \max_{x \in \mathcal{X}} F(x|\lambda) = F(\hat{x}|\lambda) \geq F(x_d|\lambda) > 0, \\ \hat{x} \in \mathcal{X}(\lambda), \quad \lambda \in \Delta.$$

The sign-definiteness of the ex post optimal objective is critical for its role as a reference, against which the goal achievement of any feasible decision can be compared.

Remark 4. The nontriviality condition (N) can be satisfied without affecting $\mathcal{X}(\cdot)$, that is, without changing any set of ex post optimal decisions. It is sufficient to consider the translated criterion $\hat{f}_i = f_i + s_i$ instead of f_i , for all $i \in \mathcal{N}$, using a suitable shift $s = (s_1, \dots, s_n) \in \mathbb{R}_{++}^n$, in which case $\hat{F}(x|\lambda) = F(x|\lambda) + c_0(\lambda) > 0$, because $c_0(\lambda) = \lambda \cdot s$ is strictly positive (as it is bounded from below by $\min\{s_1, \dots, s_n\} > 0$).

2.2. Comparative Statics

What happens to the criteria at the optimum when shifting weight from one criterion to another? At the optimum, one would naturally expect that a criterion which receives relatively more weight than before cannot decrease, whereas a criterion that receives relatively less weight cannot increase. The following result formalizes this intuition for a weight shift from one criterion to another.

Proposition 1. Let $\delta_{ij} = e_i - e_j$, where e_i and e_j denote the i -th and j -th Euclidean unit vectors, respectively. Consider $\lambda, \hat{\lambda} \in \Delta$ such that $\hat{\lambda} = \lambda + \varepsilon \delta_{ij}$ for some $\varepsilon > 0$ and some

$i, j \in \mathcal{N}$ with $i \neq j$. Then for any $(x, \hat{x}) \in \mathcal{X}(\lambda) \times \mathcal{X}(\hat{\lambda})$ it is $f_i(x) \leq f_i(\hat{x})$ and $f_j(\hat{x}) \leq f_j(x)$.

A transfer from the j -th criterion weight λ_j to the i -th criterion weight λ_i augments the optimal value of criterion i and lowers the optimal value of criterion j (at least weakly). Criteria other than i and j , whose weights remain constant but whose values are affected when decisions change, may go either way as a result of the weight shift. Similarly, the value function $F^*(\lambda)$ could go up or down, depending primarily on the difference between f_i and f_j at the optimum.

Remark 5. The conclusion of Proposition 1 can be applied multiple times. In particular, it can also be used for shifts in nonnormalized weights $w = (w_j, w_{-j}) \in \mathbb{R}_+^n \setminus \{0\}$; see Remark 2. Thus, increasing w_j is equivalent to (at most) $n - 1$ successive weight transfers in the normalized weight from the components of λ_{-j} to λ_j , resulting in an increase of the j -th criterion at the optimum.

Remark 6. For small shifts of weight from criterion j to criterion i , the value of the ex post optimal objective $F^*(\lambda) = F(x^*(\lambda)|\lambda)$ goes up (resp., down) when the corresponding score difference, $f_i(x^*(\lambda)) - f_j(x^*(\lambda))$, is positive (resp., negative), at the selection $x^*(\lambda) \in \mathcal{X}(\lambda)$.

The following result states that for a simple nonnormalized increase of the i -th weight, the value function goes up, as long as the i -th criterion is always positive at an optimum.

Lemma 1. Let $w, \hat{w} \in \mathbb{R}^n \setminus \{0\}$ be nonnormalized weights, such that $\hat{w} = w + \delta e_i$, for a given $\delta > 0$ and a given $i \in \mathcal{N}$. Provided that $f_i(x) > 0$, for any $x \in \mathcal{X}(w)$, it is $F^*(\hat{w}) > F^*(w)$.¹⁵

2.3. Performance Index

The decision maker may a priori have no knowledge about which weight $\lambda \in \Delta$ should be used to compute the weighted objective $F(\cdot|\lambda)$ in Equation (1).¹⁶ To deal with this ambiguity, the decision maker evaluates any feasible decision $x \in \mathcal{X}$ with respect to any given weight $\lambda \in \Delta$ by the performance ratio,¹⁷

$$\varphi(x|\lambda) = \frac{F(x|\lambda)}{F^*(\lambda)} \in [0, 1], \quad (3)$$

which is continuous on $\mathcal{X} \times \Delta$ and naturally bounded from above by one. Its minimum with respect to all possible weights (cf. Endnote 13),

$$\rho(x) = \min_{\lambda \in \Delta} \varphi(x|\lambda) \in [0, 1], \quad (4)$$

is called the *performance index (evaluated at x)*. The performance index measures the performance of the action x relative to all possible weighted-sum scalarizations of the decision maker's multicriteria decision problem. It

turns out that this performance criterion depends only on the maximized individual criteria,

$$f_i^* = \max_{x \in \mathcal{X}} f_i(x), \quad i \in \mathcal{N}. \quad (5)$$

By the nontriviality condition (N), it is $f_i^* = F^*(e_i) \geq F(x_d|e_i) > 0$, for all $i \in \mathcal{N}$. Thus, all maximized individual criteria are strictly positive. The next result provides an important representation of the performance index, in terms of relative goal achievement of a decision, relative to the various maximized individual criteria.

Proposition 2. The performance index is equal to

$$\rho(x) = \min_{i \in \mathcal{N}} \phi_i(x), \quad x \in \mathcal{X}, \quad (6)$$

where $\phi_i(x) = f_i(x)/f_i^* \in [0, 1]$, for all $(x, i) \in \mathcal{X} \times \mathcal{N}$.

The representation in Equation (6) expresses the performance index as the minimum of the criterion-specific performance ratios ϕ_i . Hence, to compute $\rho(x)$ as the minimum of $\varphi(x|\lambda)$ over all weights $\lambda \in \Delta$, it is sufficient to restrict attention to the Euclidean unit vectors $e_i \in \Delta$, as $\phi_i(\cdot) = F(\cdot|e_i)/F^*(e_i)$, for all $i \in \mathcal{N}$. This simplification arises because $\varphi(x|\cdot)$ is quasiconcave for fixed x , implying that its minimum over Δ is attained at a vertex. This highlights a “perfect complementarity” among the criterion-specific performance ratios in the determination of the performance index.

Remark 7. The idea of perfect complementarity, discussed by Cournot (1838) and Edgeworth (1897), describes elements contributing to a common goal in fixed proportions. This is equivalent to evaluating the criterion (i.e., the performance index) using the Leontief production function,¹⁸ that is, taking the minimum among the inputs ϕ_i , resulting in $\rho(x) = \min_{i \in \mathcal{N}} \phi_i(x)$, for all $x \in \mathcal{X}$, as in Proposition 2.

Example 1. Let $\mathcal{X} \subset \mathbb{R}_+^m$ be a compact action set with $\mathcal{X} \cap \mathbb{R}_+^m \neq \emptyset$, where $m \geq 2$ is an integer. Consider a decision problem with $n = 2$ criteria (containing an egalitarian and a utilitarian evaluation),

$$f_1(x) = \min\{x_1, \dots, x_m\} \text{ and } f_2(x) = (1/m)(x_1 + \dots + x_m),$$

for all $x = (x_1, \dots, x_m)$, with maximized values $f_i^* = \max_{x \in \mathcal{X}} f_i(x) > 0$, for $i \in \mathcal{N}$. The corresponding weighted objective becomes

$$F(x|\ell) = (1 - \ell)f_1(x) + \ell f_2(x), \quad \ell \in [0, 1],$$

where, for simplicity, the parameter ℓ replaces the normalized weight $\lambda = (1 - \ell, \ell)$. The optimal value, $F^*(\ell) = \max_{x \in \mathcal{X}} F(x|\ell)$, can be obtained by solving a two-stage maximization problem,

$$F^*(\ell) = \max_{t \in \mathcal{T}} \{(1 - \ell)t + \ell \mu(t)\}, \quad \ell \in [0, 1].$$

Here $\mathcal{T} = [t_1, t_2]$ is a compact interval, with $0 < t_1 \leq t_2 < \infty$, and the best average coordinate (subject to all

strictly improves on at least one criterion while weakly improving on all other criteria (over their values achieved at x). The corresponding preference relation $>$ on $\mathcal{X}$ is defined by²⁰

$$x' > x \Leftrightarrow (\exists i \in \mathcal{N} : f_i(x') > f_i(x), f(x') \geq f(x)), \quad (9)$$

for all $x, x' \in \mathcal{X}$. The resulting set of *efficient* (or *Pareto-optimal*) decisions is

$$\mathcal{P} = \{x \in \mathcal{X} : (f(x) \leq f(x') \Rightarrow f(x) = f(x')), x' \in \mathcal{X}\}. \quad (10)$$

It is easy to see that a pseudo-robust decision $\hat{x} \in \Psi$ is not necessarily efficient, in the sense that it may be possible to find a different decision which strictly improves on at least one criterion-specific performance ratio while maintaining the optimal performance index achieved by $\hat{x}$.

Example 3. Following up on our analysis in Example 1 and Example 2, consider the (nonconvex, compact) action set $\mathcal{X} = [0, 1]^m \setminus (1/2, 1]^m$, for some integer $m \geq 2$. Because $\mu(t) = 1 - 1/(2m)$, for all $t \in \mathcal{T} = [t_1, t_2]$, where $t_1 \leq t_2 = f_1^* = 1/2$ and $f_2^* = \mu(t_1)$, the optimality condition $\hat{t}/t_2 = \mu(\hat{t})/\mu(t_1) = 1$ yields $\hat{t} = 1/2$. Hence, the optimal performance index attains its maximum possible value, $\rho^* = 1$. This makes sense, because it is feasible (in $\mathcal{X}$) to attain simultaneously the highest possible minimum coordinate and the highest possible average coordinate, regardless of the weights assigned to the two objectives. Note also that the set of pseudo-robust actions,

$$\begin{aligned} \Psi &= \{x \in \mathcal{X} : \min\{x_1, \dots, x_m\} = \hat{t}/\rho^*\} \\ &= \{x \in [1/2, 1]^m : \exists i \in \mathcal{N} \text{ s.t. } x_i = 1/2\}, \end{aligned}$$

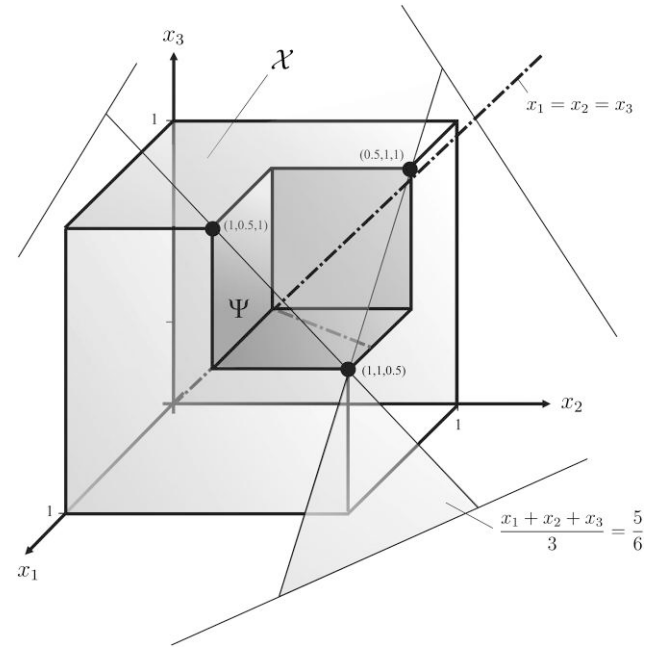
has a continuum of elements. For instance, compared with the pseudo-robust action $(1, \dots, 1)/2$, all other actions in Ψ are more efficient, particularly those in the set of Pareto-optimal decisions, $\mathcal{P} = \{(1, \dots, 1) - (e_i/2) : i \in \mathcal{N}\}$, which is a subset of Ψ . For $m = 2$, we obtain that $\mathcal{P} = \{(0.5, 1), (1, 0.5)\}$. Figure 3 depicts the situation, including the set of Pareto-optimal decisions, for $m = 3$.

Remark 8. It is well known that for any weight vector λ with strictly positive components (ensuring all criteria are considered) an ex post optimal decision must also be efficient (see, e.g., Ehrgott 2010, proposition 3.9, p. 71). That is, for any $\lambda \in \text{int}(\Delta)$, the corresponding optimal decision set satisfies $\mathcal{X}(\lambda) \subset \mathcal{P}$.

3.3. Robust Decision Set

Requiring efficiency in addition to pseudo-robustness is important, because, by avoiding unnecessary shortfalls, it can only improve the weighted objective in Equation (1)—at least weakly. A decision $\hat{x}^* \in \mathcal{X}$ is called *robust* if

Figure 3. Nonconvex Action Set in Example 3, for $m = 3$



Note. The set of Pareto-optimal decisions is $\mathcal{P} = \{(1, 1, 0.5), (1, 0.5, 1), (0.5, 1, 1)\}$.

it is pseudo-robust and efficient. Thus, our goal becomes to examine the properties of the *robust decision set*,

$$\mathcal{R} = \Psi \cap \mathcal{P}, \quad (11)$$

which contains all available robust decisions in $\mathcal{X}$, and to then determine a direct method for the computation of a robust decision (or an arbitrarily close approximation thereof).

Example 4. Following up on our analysis in Example 1 and Example 2, consider the (convex, compact) action set $\mathcal{X} = \{x \in \mathbb{R}^m : \|x\|_2 \leq 1\}$, which is equal to the unit ball in the standard Euclidean distance, where the dimension of the underlying space is given by some integer $m \geq 2$. By direct computation, $\mu(t) = 1/\sqrt{m}$ for all $t \in \mathcal{T} = [t_1, t_2]$, and $t_1 = f_1^* = 1/\sqrt{m} = f_2^* = \mu(t_2) = t_2$. As in Example 3, the balancedness condition (8) yields $\hat{t} = t_2$, and thus an optimal performance index of $\rho^* = 1$. The set of pseudo-robust actions,

$$\begin{aligned} \Psi &= \{x \in \mathcal{X} : \min\{x_1, \dots, x_m\} = \hat{t}/\rho^*\} \\ &= \{(1/\sqrt{m}, \dots, 1/\sqrt{m})\}, \end{aligned}$$

is a singleton. On the other hand, the set of Pareto-optimal actions,

$$\mathcal{P} = \{x \in \mathbb{R}_+^m : \|x\|_2 = 1\},$$

contains a continuum of elements, including the only pseudo-robust decision, $(1, \dots, 1)/\sqrt{m}$.

In Example 3, we examined a problem where $\mathcal{P} \subsetneq \Psi$ (so $\mathcal{R} = \mathcal{P}$), whereas in Example 4, for the same

multicriteria decision problem with a different action set, it was $\Psi \subsetneq \mathcal{P}$ (so $\mathcal{R} = \Psi$). Neither of these two extremes might apply, in which case $\mathcal{R} \notin \{\Psi, \mathcal{P}\}$, as illustrated next.

Example 5. Mixing and matching features from Example 3 and Example 4, let us consider the (nonconvex, compact) action set $\mathcal{X} = \{x \in \mathbb{R}^2 : \|x\|_2 \leq 1\} \setminus (1/2, 1)^2$ in the Euclidean plane. The corresponding set of Pareto-optimal actions is given by the intersection of the unit circle with both the action set $\mathcal{X}$ and the positive quadrant $\mathbb{R}_+^2$, so $\mathcal{P} = \{x \in \mathcal{X} \cap \mathbb{R}_+^2 : \|x\|_2 = 1\}$. Meanwhile, the set of pseudo-robust actions is $\Psi = \{x \in \mathcal{X} : \min\{x_1, x_2\} = 1/2\}$. Thus, by Equation (11) it is $\mathcal{R} = \Psi \cap \mathcal{P} = \{(1/2, \sqrt{3}/2), (\sqrt{3}/2, 1/2)\}$, which in this setting means $\mathcal{R} \notin \{\Psi, \mathcal{P}\}$.

The existence of robust actions is ensured by the next result, together with the fact that the robust decision set $\mathcal{R}$ must be closed and bounded (i.e., compact).

Lemma 2. *The robust decision set $\mathcal{R}$ is nonempty and compact.*

The proof of Lemma 2 starts by noting that the set $\mathcal{P}$ of efficient actions is compact because it must be bounded (by the boundedness of the encompassing action set $\mathcal{X}$) and closed (by the continuity of the criteria). One can then construct a sequence of pseudo-robust actions in the (nonempty) compact set Ψ which might be inefficient (or else $\mathcal{R} \neq \emptyset$ holds true immediately). But each inefficient pseudo-robust action suggests the existence of a more efficient action, which incidentally must also be pseudo-robust. Given that Ψ is compact, the Bolzano-Weierstrass theorem (Berge 1963, p. 67) then guarantees the existence of a converging subsequence of pseudo-robust actions, with a limit that must be a feasible decision which is both pseudo-robust and efficient, so $\mathcal{R} \neq \emptyset$. Compactness of the robust decision set then follows, as it has to be both closed and bounded.

Lemma 3. *The optimal performance index ρ^* is such that $\max \rho(\mathcal{P}) = \rho^* = \max \rho(\mathcal{X})$ and $\rho(\mathcal{R}) = \rho(\Psi) = \{\rho^*\}$.*

The preceding result (re)states the fact that the decision maker can restrict attention to efficient actions when maximizing the performance index, meaning that there always exists an efficient action which attains the optimal performance index ρ^* ; this action is by definition robust. Conversely, any robust action necessarily achieves a performance index of ρ^* , which is quite straightforward in light of both Lemma 2 and the definition of the robust decision set in Equation (11).

3.4. Robust Decisions: Computation

How can one determine a (relatively) robust decision? By Lemma 3 we can limit attention to maximizing the performance index over all efficient decisions in our search for robust decisions. Thus, to be guaranteed an

efficient decision which is also approximately pseudo-robust, by virtue of Proposition 2 we introduce the ε -augmented performance index,

$$\Phi_\varepsilon(x) = (1 - \varepsilon) \min_{i \in \mathcal{N}} \phi_i(x) + (\varepsilon/n) \sum_{i \in \mathcal{N}} \phi_i(x), \quad (x, \varepsilon) \in \mathcal{X} \times [0, 1], \quad (12)$$

so that $\Phi_0 = \rho = \min_{i \in \mathcal{N}} \phi_i$, and $\Phi_1 = (1/n) \sum_{i \in \mathcal{N}} \phi_i$. Because $\Phi_\varepsilon(\cdot)$ is a continuous function, its maximizer, referred to as the set of ε -robust actions,

$$\mathcal{R}_\varepsilon = \arg \max_{x \in \mathcal{X}} \Phi_\varepsilon(x), \quad \varepsilon \in [0, 1], \quad (13)$$

is nonempty and compact, again by virtue of the extreme-value theorem and the maximum theorem. It is useful, for our further analysis, to shift the perspective from the available decisions in $\mathcal{X}$ to their respective consequences (or “outcomes”) in terms of their robustness performance, given the prevailing weight ambiguity.

Remark 9. The (nonempty, compact) outcome space,

$$\mathcal{Y} = \{y \in [0, 1]^n : y_i = \phi_i(x), (x, i) \in \mathcal{X} \times \mathcal{N}\}, \quad (14)$$

contains the vectors y of criterion-specific performance ratios, achieved by the available decisions x , so $\mathcal{Y} = \phi(\mathcal{X})$, where $\phi = (\phi_1, \dots, \phi_n)$. As several feasible decisions might result in the same score vector, the mapping $\phi : \mathcal{X} \rightarrow \mathcal{Y}$ may not be one-to-one. Consider now the ε -augmented performance index in the outcome space,

$$\hat{\Phi}_\varepsilon(y) = (1 - \varepsilon) \min_{i \in \mathcal{N}} y_i + (\varepsilon/n) \sum_{i \in \mathcal{N}} y_i, \quad (y, \varepsilon) \in \mathcal{Y} \times \mathbb{R}_+.$$

Using ideas from Example 1, maximization of this weighted objective yields

$$\hat{\Phi}_\varepsilon^* = \max_{y \in \mathcal{Y}} \hat{\Phi}_\varepsilon(y) = \max_{t \in \mathcal{T}} \{(1 - \varepsilon)t + \varepsilon \hat{\mu}(t)\}, \quad \varepsilon \in [0, 1],$$

where $\mathcal{T} = [t_1, t_2]$ is a suitable compact interval, and

$$\hat{\mu}(t) = \max \left\{ (1/n) \sum_{i \in \mathcal{N}} y_i : t \leq y_i, (y, i) \in \mathcal{Y} \times \mathcal{N} \right\}, \quad t \in \mathcal{T},$$

denotes the *average (criterion-specific) performance ratio*. The interval boundaries t_1, t_2 of $\mathcal{T}$, with $0 < t_1 \leq t_2 \leq 1$, are given by

$$t_1 = \min_{y \in \mathcal{Y}} \min_{i \in \mathcal{N}} y_i \quad \text{and} \quad t_2 = \max_{y \in \mathcal{Y}} \min_{i \in \mathcal{N}} y_i.$$

In addition, one can easily verify that

$$\begin{aligned} \hat{\Phi}_0^* &= t_2 = \rho^* = \max \rho(\mathcal{X}) \quad \text{and} \\ \hat{\Phi}_1^* &= \hat{\mu}(t_1) = (1/n) \max_{y \in \mathcal{Y}} \sum_{i \in \mathcal{N}} y_i \geq t_2. \end{aligned}$$

Proposition 1 implies that any selection $t_\varepsilon \in \mathcal{T}_\varepsilon = \arg \max_{t \in \mathcal{T}} \{(1 - \varepsilon)t + \varepsilon \hat{\mu}(t)\}$ must be decreasing in ε , and similarly, $\hat{\mu}(t_\varepsilon)$ must be increasing in ε , at least

weakly. The latter also follows directly from the monotonicity of t_ε , because $\hat{\mu}(t)$ must be nonincreasing in t .

Lemma 4. Let $\varepsilon \in (0, 1]$. (i) $\mathcal{R}_0 = \Psi$. (ii) $\mathcal{R}_\varepsilon \subset \mathcal{P}$. (iii) If $x \in \mathcal{R}_\varepsilon \setminus \Psi$ and $\hat{x} \in \Psi$, then there exists $\varepsilon' \in (0, \varepsilon)$ such that $\Phi_{\hat{\varepsilon}}(x) < \Phi_{\hat{\varepsilon}}(\hat{x})$, for all $\hat{\varepsilon} \in [0, \varepsilon']$.

Part (i) of the preceding result notes that without ε -augmentation (i.e., for $\varepsilon = 0$) the weighted objective Φ_ε specializes to the performance index (via Proposition 2), the maximization of which produces the set of pseudo-robust decisions, as in Equation (7). Part (ii) stipulates that whenever the ε -augmentation is *nontrivial* (i.e., for $\varepsilon > 0$), maximization yields efficient actions. Finally, part (iii) means that if a decision x maximizes the ε -augmented performance index in Equation (12), for some nontrivial $\varepsilon \in (0, 1]$, without being pseudo-robust, then any pseudo-robust decision $\hat{x} \in \Psi$ would strictly improve upon x in terms of any $\hat{\varepsilon}$ -augmented performance index, as long as $\hat{\varepsilon}$ (smaller than ε) lies in a sufficiently small right-neighborhood of the origin.

We are now able to establish a cornerstone property for the practice of robust multicriteria optimization, in the sense that a robust decision (i.e., an element of $\mathcal{R}$) can be obtained as a lower limit of the set of ε -robust actions, for $\varepsilon \rightarrow 0^+$.²¹ In Section 3.5, we then show that $\varepsilon > 0$ can be chosen so as to guarantee an approximation of the optimal performance ρ^* up to any desired precision.

Proposition 3. Let $\mathcal{Q} = \underline{\text{Lim}}_{\varepsilon \rightarrow 0^+} \mathcal{R}_\varepsilon$. Then $\mathcal{Q} \subset \mathcal{R}$ and $\mathcal{Q} \neq \emptyset$.

The following example shows that it is possible that $\mathcal{Q} \subsetneq \mathcal{R}$. In other words, some points in the robust decision set might not be reached using the proposed approximation procedure.

Example 6. Consider a (nonconvex, compact) action set as shown in Figure 4(a), specified by

$$\mathcal{X} = \{x \in [0, 1] \times \mathbb{R}_+ : x_2 \leq 2 + \sqrt{1 - x_1}\} \cup \{(2, \frac{3}{2})\}.$$

There are $n = 2$ criteria that simply measure the coordinate achievement, so $f_i(x) = x_i$, for all $(x, i) \in \mathcal{X} \times \mathcal{N}$. The set of efficient actions is

$$\mathcal{P} = \{x \in [0, 1] \times \mathbb{R}_+ : x_2 = 2 + \sqrt{1 - x_1}\} \cup \{(2, \frac{3}{2})\},$$

whereas the set of pseudo-robust actions is given by $\Psi = \{(1, x_2) : 1 \leq x_2 \leq 2\} \cup \{(2, \frac{3}{2})\}$. By its definition in Equation (11) the set of robust actions can be obtained as the intersection of the preceding sets:

$$\mathcal{R} = \Psi \cap \mathcal{P} = \{(1, 2), (2, \frac{3}{2})\}.$$

At this point, let us consider the ε -augmented performance index in Equation (12). By Lemma 4 we have $\mathcal{R}_0 = \Psi$, and for $\varepsilon \in (0, 1]$ the maximizer of Φ_ε is efficient,

so

$$\begin{aligned} \Phi_\varepsilon^* &= \max_{x \in \mathcal{P}} \Phi_\varepsilon(x) \\ &= \max \left\{ \max_{x_1 \in [0, 1]} \left\{ \frac{(1 - \varepsilon)x_1}{2} + \frac{\varepsilon}{2} \left(\frac{x_1}{2} + \frac{2 + \sqrt{1 - x_1}}{3} \right) \right\}, \right. \\ &\quad \left. \Phi_\varepsilon(2, \frac{3}{2}) \right\}, \end{aligned}$$

where $\phi_i = f_i/f_i^*$, for $i \in \mathcal{N}$, with $(f_1^*, f_2^*) = (2, 3)$, and $\Phi_\varepsilon(2, \frac{3}{2}) = (2 + \varepsilon)/4$. By direct computation one finds that the maximizer of the (nontrivially) ε -augmented performance index,²²

$$\mathcal{R}_\varepsilon = \left\{ \left(2, \frac{3}{2} \right) \right\}, \quad \varepsilon \in (0, 1],$$

is a singleton. Proposition 3 then guarantees that the lower limit of this maximizer is robust:

$$\mathcal{Q} = \underline{\text{Lim}}_{\varepsilon \rightarrow 0^+} \mathcal{R}_\varepsilon = \left\{ \left(2, \frac{3}{2} \right) \right\} \subsetneq \left\{ (1, 2), \left(2, \frac{3}{2} \right) \right\} = \mathcal{R}.$$

The fact that $(2, \frac{3}{2})$ is an isolated point (cf. Endnote 12) is not important, because one could easily connect it to $\mathcal{X}$ by adding points to the action set that are always suboptimal.²³

3.5. Approximation Error of ε -Robust Decisions

Maximizing the ε -augmented performance index Φ_ε in Equation (12) enables us to approximate a robust decision to an arbitrary prespecified precision, as a function of $\varepsilon \in (0, 1]$. The quality of any approximate decision $x_\varepsilon \in \mathcal{R}_\varepsilon$ in Equation (11) is gauged by its *approximation error*,

$$\psi_\varepsilon = \rho^* - \rho_\varepsilon \in [0, 1], \quad \varepsilon \in [0, 1], \quad (15)$$

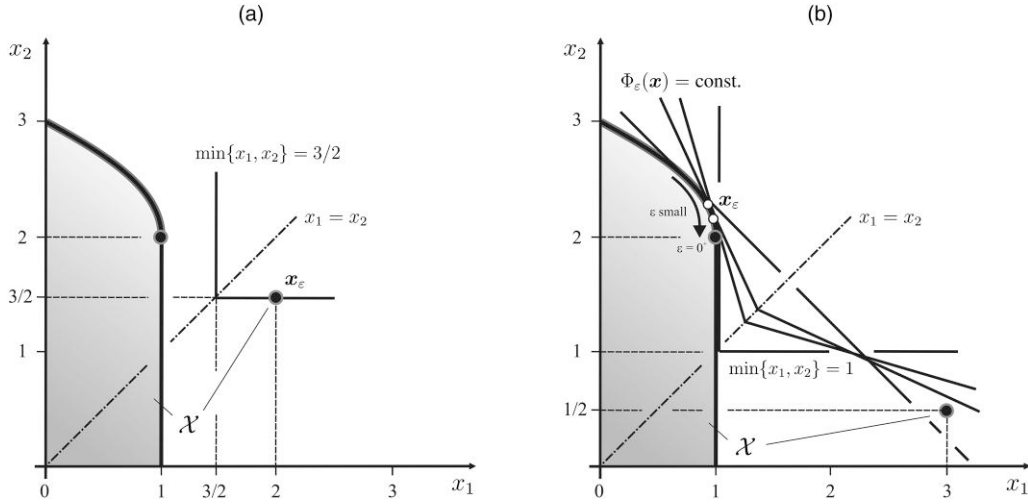
that is, the difference between the optimal performance index, $\rho^* = \max \rho(\mathcal{X})$, and the achieved performance index, $\rho_\varepsilon = \rho(x_\varepsilon)$. This quantity measures the loss in robustness from using the ε -approximation rather than the true robust decision. For the remainder of this discussion, we keep the selection $x_\varepsilon \in \mathcal{R}_\varepsilon$ fixed, for all $\varepsilon \in [0, 1]$. If we denote by $\mu_\varepsilon = \mu(x_\varepsilon) = (1/n) \sum_{i \in \mathcal{N}} \phi_i(x_\varepsilon)$ the average performance ratio achieved by x_ε , the *optimal ε -augmented performance index* becomes

$$\Phi_\varepsilon^* = \max \Phi_\varepsilon(\mathcal{X}) = (1 - \varepsilon)\rho_\varepsilon + \varepsilon\mu_\varepsilon, \quad \varepsilon \in [0, 1]. \quad (16)$$

The following result guarantees continuity, as well as first- and second-order monotonicity of the optimal ε -augmented performance index.

Lemma 5. Φ_ε^* is continuous, increasing, and convex in $\varepsilon \in [0, 1]$.

The proof of the first-order monotonicity uses the fact that $\mu_\varepsilon \geq \rho_\varepsilon$ (for all $\varepsilon \in [0, 1]$), together with natural properties of an optimal solution to Equation (13) in order to establish that Φ_ε^* in Equation (16) must increase

Figure 4. Action Sets in Examples 6 and 7

Notes. (a) Action set in Example 6 with $\lim_{\varepsilon \rightarrow 0^+} \mathcal{R}_\varepsilon = \{(2, \frac{3}{2})\} \subsetneq \mathcal{R} = \{(1, 2), (2, \frac{3}{2})\}$. (b) Action set in Example 7 with $\lim_{\varepsilon \rightarrow 0^+} \mathcal{R}_\varepsilon = \{(1, 2)\} = \mathcal{R}$.

in ε (at least weakly). That the optimal ε -augmented performance index also exhibits a second-order monotonicity means that tightening the approximation parameter further and further leads to progressively slower decreases of Φ_ε^* toward $\Phi_0^* = \rho^*$ (as $\varepsilon \rightarrow 0^+$). The convexity of the optimal ε -augmented performance index also implies its smoothness (almost everywhere), as pointed out next.

Remark 10. By the Rademacher theorem (Villani 2008, theorem 10.8, p. 222), the convexity of Φ_ε^* in $\varepsilon \in [0, 1]$, which implies Lipschitz continuity, guarantees that the optimal ε -augmented performance index is differentiable almost everywhere (a.e.) on $[0, 1]$. The Alexandrov theorem (Villani 2008, theorem 14.25, p. 402) goes even further by establishing its second-order differentiability a.e., in the sense of having a Taylor expansion with a smaller-than-quadratic local error at almost every point $\varepsilon \in [0, 1]$. By the envelope theorem (see, e.g., Mas-Colell et al. 1995, theorem M.L.1, pp. 965–966), at points of differentiability one therefore obtains:

$$\frac{d\Phi_\varepsilon^*}{d\varepsilon} = \mu_\varepsilon - \rho_\varepsilon \geq 0, \quad \forall \varepsilon \in [0, 1].$$

In other words, the gradient of the optimal ε -augmented performance index is (a.e.) equal to the difference between the average performance ratio (μ_ε) and the minimum performance ratio (ρ_ε) attained by the approximately robust decision x_ε .

The maximized weighted objective in Equation (16) is a convex combination of ρ_ε and μ_ε ; the latter both exhibit “natural” comparative statics as implied by the reweighting result in Proposition 1.

Lemma 6. (i) ρ_ε is decreasing in $\varepsilon \in [0, 1]$. (ii) $\rho^* \geq \rho_\varepsilon$ for all $\varepsilon \in [0, 1]$. (iii) $\lim_{\varepsilon \rightarrow 0^+} \rho_\varepsilon = \rho^*$. (iv) μ_ε is increasing in $\varepsilon \in [0, 1]$. (v) $\mu_\varepsilon \geq \rho^*$ for all $\varepsilon \in [0, 1]$.

Parts (i) and (iv) of Lemma 6 assert that decreasing the augmentation parameter ε can only decrease the average criterion-specific performance ratio μ_ε and at the same time increase the performance index ρ_ε , where both are achieved at a given ε -robust selection x_ε . Meanwhile, by parts (ii) and (v) the optimal performance ratio $\rho^* = \max \rho(\mathcal{X})$ is bracketed by these two values, in the sense that $\rho_\varepsilon \leq \rho^* \leq \mu_\varepsilon$ for all $\varepsilon \in [0, 1]$. Finally, parts (i) and (iii) establish the monotone convergence of ρ_ε to ρ^* , as $\varepsilon \rightarrow 0^+$.

The following result provides a *performance guarantee* for any ε -approximation of our robust multicriteria decision problem, alluded to at the outset of our discussion.

Proposition 4. Fix any $\delta > 0$. If $\varepsilon \leq \delta/\mu_1$, then $\psi_\varepsilon \leq \delta$.

A special case of the preceding result is that for any $\delta > 0$ the approximation error ψ_{δ/μ_1} cannot exceed δ . In other words, any ε -robust decision x_ε , for $\varepsilon = \delta/\mu_1$, attains a performance index $\rho_\varepsilon \in [\rho^* - \delta, \rho^*]$, where ρ^* is the optimal performance index. The following example applies this result in an already familiar context.

Example 7. Somewhat similar to Example 6, we consider the (nonconvex, compact) action set

$$\mathcal{X} = \{x \in [0, 1] \times \mathbb{R}_+ : x_2 \leq 2 + \sqrt{1 - x_1}\} \cup \{(3, \frac{1}{2})\},$$

together with the same $n = 2$ criteria $f_i(x) = x_i$, for all $(x, i) \in \mathcal{X} \times \mathcal{N}$. Consequently, the set of efficient actions is $\mathcal{P} = \{x \in [0, 1] \times \mathbb{R}_+ : x_2 = 2 + \sqrt{1 - x_1}\} \cup \{(3, \frac{1}{2})\}$, and the set of pseudo-robust actions is $\Psi = \{(1, x_2) : 1 \leq x_2 \leq 2\}$. Thus, by Equation (11) the set of robust actions amounts to $\mathcal{R} = \Psi \cap \mathcal{P} = \{(1, 2)\}$, which is a

singleton. By Lemma 4 the ε -augmented performance index Φ_ε in Equation (12) is efficient (i.e., $\mathcal{R}_\varepsilon \subset \mathcal{P}$), for all $\varepsilon \in (0, 1]$, and $\mathcal{R}_0 = \Psi$. In particular, the optimal value of the ε -augmented performance index is

$$\begin{aligned}\Phi_\varepsilon^* &= \max_{x \in \mathcal{P}} \Phi_\varepsilon(x) \\ &= \max \left\{ \max_{x_1 \in [0, 1]} \left\{ \frac{(1-\varepsilon)x_1}{3} + \frac{\varepsilon}{2} \left(\frac{x_1}{3} + \frac{2 + \sqrt{1-x_1}}{3} \right) \right\}, \right. \\ &\quad \left. \Phi_\varepsilon(3, \tfrac{1}{2}) \right\},\end{aligned}$$

where $\phi_i = f_i/f_i^*$, for $i \in \mathcal{N}$, with $(f_1^*, f_2^*) = (3, 3)$, and $\Phi_\varepsilon(3, \frac{1}{2}) = (2 + 5\varepsilon)/12$. For sufficiently small ε , the maximizer,

$$\begin{aligned}\mathcal{R}_\varepsilon &= \left\{ \left(1 - \left(\frac{\varepsilon/2}{2-\varepsilon} \right)^2, 2 + \frac{\varepsilon/2}{2-\varepsilon} \right) \right\}, \\ 0 < \varepsilon < \frac{4}{7} \left(2 - \frac{1}{\sqrt{2}} \right) &\approx 0.7388,\end{aligned}$$

is single-valued, achieving $\Phi_\varepsilon^* = (1/24)(16 - 3\varepsilon^2)/(2 - \varepsilon)$.²⁴ By Proposition 3 the lower limit of this maximizer for $\varepsilon \rightarrow 0^+$ is robust, and in our case:

$$\mathcal{Q} = \lim_{\varepsilon \rightarrow 0^+} \mathcal{R}_\varepsilon = \{(1, 2)\} = \mathcal{R};$$

see Figure 4(b). Note also that $\mathcal{R}_1 = \{(3, \frac{1}{2})\}$, which in turn yields the average criterion-specific performance ratio achieved by the decision $x_\varepsilon \in \mathcal{R}_\varepsilon$, for $\varepsilon = 1$:

$$\mu_1 = \frac{1}{2} \left(\frac{x_1}{f_1^*} + \frac{x_2}{f_2^*} \right) \Big|_{(x_1, x_2) = (3, \frac{1}{2})} = \frac{7}{12} \approx 0.5833.$$

Thus, to guarantee that the approximation error ψ_ε in Equation (15) cannot exceed $\delta = 5\%$, it is by Proposition 4 enough to find an action x_ε that maximizes the ε -augmented performance index Φ_ε in Equation (12) for some $\varepsilon \in (0, \delta/\mu_1] \approx (0, 8.6\%]$. Incidentally, for $\varepsilon = 8\%$, we find $\rho_\varepsilon \approx 33.32\%$, so that the realized approximation error of $\psi_\varepsilon \approx 0.01\%$ (with respect to $\rho^* = 1/3$) is actually much smaller than the prespecified 5% approximation-error bound.

It is also possible to derive a priori performance estimates without knowledge of the optimal performance index (ρ^*) just by solving the ε -approximation problem in Equation (13) for some admissible ε . Indeed, Lemma 6(ii) and Lemma 5 together imply that

$$\rho_\varepsilon \leq \rho^* \leq \Phi_\varepsilon^*, \quad \varepsilon \in [0, 1]. \quad (17)$$

For any $\varepsilon \in [0, 1]$, let us now consider the midpoint estimator,

$$\begin{aligned}\hat{\rho}_\varepsilon &= \frac{\rho_\varepsilon + \Phi_\varepsilon^*}{2} = \rho_\varepsilon + \varepsilon \frac{\mu_\varepsilon - \rho_\varepsilon}{2} \\ &= \rho(x_\varepsilon) + \varepsilon \frac{\mu(x_\varepsilon) - \rho(x_\varepsilon)}{2}, \quad \varepsilon \in [0, 1].\end{aligned} \quad (18)$$

The latter means that $\hat{\rho}_\varepsilon$ effectively approximates ρ^* , for $\varepsilon \rightarrow 0^+$.

Lemma 7. $|\hat{\rho}_\varepsilon - \rho^*| \leq \varepsilon(\mu_\varepsilon - \rho_\varepsilon)/2$, for all $\varepsilon \in [0, 1]$.

Because $\mu_\varepsilon - \rho_\varepsilon \in [0, 1]$, a somewhat simpler (though generally less precise) approximation inequality than the one given in the preceding Lemma 7 is $|\hat{\rho}_\varepsilon - \rho^*| \leq \varepsilon/2$, for all $\varepsilon \in [0, 1]$.

Remark 11. For a robust decision $\hat{x} \in \mathcal{R}$ let $\hat{y} = (f_1(\hat{x}), \dots, f_n(\hat{x}))$; in addition, let $y^* = (f_1^*, \dots, f_n^*)$ be the utopian point in the outcome space. Because by construction $\hat{y}_i \geq \rho^* y_i^*$, for all $i \in \mathcal{N}$, it is

$$\rho^* \geq \underline{\rho} = \sup\{r \in [0, 1] : r y^* \in \mathcal{Y}\},$$

where $\sup \emptyset = 0$. In the *balanced case*, where $f_i(\hat{x}) = \rho^* f_i^*$, for all $i \in \mathcal{N}$, the lower bound becomes tight (i.e., $\rho^* = \underline{\rho}$). However, in general $\underline{\rho}$ may not be very reliable (e.g., generically worse than the performance index $\rho_d = \rho(x_d)$, attained by the default decision x_d).

3.6. Robust Weights

3.6.1. General Case. An interesting and useful byproduct of a robust decision $\hat{x} \in \mathcal{R}$ is its associated *robust (normalized) weight*,

$$\hat{\lambda} = \left(\frac{\hat{h}}{f_1(\hat{x})}, \dots, \frac{\hat{h}}{f_n(\hat{x})} \right) \in \Delta, \quad (19)$$

where the positive constant

$$\hat{h} = \left(\sum_{i \in \mathcal{N}} \frac{1}{f_i(\hat{x})} \right)^{-1} \quad (20)$$

denotes the harmonic mean of the criterion scores at the robust decision. By construction one obtains that the contribution to the robustly weighted objective $F(\cdot | \hat{\lambda})$, evaluated at the robust decision $\hat{x}$, is equal for all criteria, in the sense that

$$\max_{x \in \mathcal{X}} \min_{i \in \mathcal{N}} \hat{\lambda}_i f_i(x) = \hat{h} = \hat{\lambda}_i f_i(\hat{x}), \quad i \in \mathcal{N}. \quad (21)$$

When the action set is convex, then the chosen robust action naturally also maximizes the robustly weighted criterion, that is, $\hat{x} \in \mathcal{X}(\hat{\lambda})$. One can think of the robust weight as an endogenous belief that can be used to evaluate the expected criterion achievement. It defines the tradeoffs compatible with the robust decision, where the latter was found while remaining agnostic over all possible weights (in Δ).

3.6.2. Balanced Case. In the case where $f_i(\hat{x}) = \rho^* f_i^*$, for all $i \in \mathcal{N}$ (cf. Remark 11), we have

$$\hat{\lambda} = \left(\frac{h^*}{f_1^*}, \dots, \frac{h^*}{f_n^*} \right), \quad (19')$$

where

$$h^* = \left(\sum_{i \in \mathcal{N}} \frac{1}{f_i^*} \right)^{-1} \quad (20')$$

is the harmonic mean of the maximum criterion scores.

The balancedness condition then becomes

$$\rho^* = \frac{f_i(\hat{x})}{f_i^*} = \frac{\hat{h}}{h^*}, \quad i \in \mathcal{N}.$$

In the balanced case, the robust weight can therefore be determined based on the maximized individual criteria alone. If all maximized individual criteria are equal, the robust weight becomes uniform. Indeed, $f_i^* = c > 0$, for all $i \in \mathcal{N}$, implies that $h^* = c/n$, so $\hat{\lambda} = (1/n, \dots, 1/n)$.²⁵

Example 8. In Example 1, we determined the value function $F^*(\ell) = (1 - \ell)\hat{t} + \ell\mu(\hat{t})$, so that by Equations (19') and (20') in this balanced case (as established in Example 2):

$$\hat{\lambda} = (1 - \hat{\ell}, \hat{\ell}) = \left(\frac{\mu(\hat{t})}{\hat{t} + \mu(\hat{t})}, \frac{\hat{t}}{\hat{t} + \mu(\hat{t})} \right) = \left(\frac{f_2^*}{f_1^* + f_2^*}, \frac{f_1^*}{f_1^* + f_2^*} \right).$$

Generically, it is $F^*(\hat{\ell}) \neq F(\hat{x}|\hat{\ell})$, even when the decision set is convex. To see this, let $m = n = 2$ and consider the convex domain, $\mathcal{X} = \{x \in \mathbb{R}_+^2 : v_1x_1 + v_2x_2 \leq 1\}$, where the constants v_1, v_2 are such that $0 < v_1 < v_2 < \infty$. Then $\mu(t) = (1 - v_2t)/v_1 + t$, so that $\mu'(t) = 1 - (v_2/v_1) < 0$. Meanwhile, $F(t|\ell) = (1 - \ell)t + \ell\mu(t)$, for $t \in [0, f_1^*] = [0, 1/(v_1 + v_2)]$, and

$$F'(t|\ell) = (1 - \ell) + \ell\mu'(t) = 1 - \frac{\ell v_2}{v_1} \geq 0 \Leftrightarrow \ell \leq \frac{v_1}{v_2}.$$

Therefore,

$$F^*(\ell) = F(t^*(\ell)|\ell) = (\ell/v_1) + \max\{0, 1 - (\ell v_2/v_1)\}f_1^*,$$

where

$$t^*(\ell) = \begin{cases} \{f_1^*\}, & \text{if } \ell < v_1/v_2, \\ [0, f_1^*], & \text{if } \ell = v_1/v_2, \\ \{0\}, & \text{if } \ell > v_1/v_2, \end{cases}$$

and $f_1^* = 1/(v_1 + v_2)$. Because $f_2^* = 1/v_1$, we find $v_1/v_2 = f_1^*/(f_2^* - f_1^*)$. By virtue of the fact that $\hat{\ell} = f_1^*/(f_1^* + f_2^*) < v_1/v_2$, it is $t^*(\hat{\ell}) = f_1^*$. The balancedness condition,

$$\frac{\hat{t}}{f_1^*} = \frac{\mu(\hat{t})}{f_2^*},$$

is equivalent to $\hat{t}/f_1^* = 1 - (\hat{t}/f_1^*)$, so $\hat{t} = f_1^*/2$. Hence, we find that $\hat{t} \neq t^*(\hat{\ell})$. Let us check the corresponding performance index. Indeed,

$$\rho(\hat{t}) = 0.5 > 0 = \min\left\{\frac{f_1^*}{f_1^*}, 1 - \frac{f_1^*}{f_1^*}\right\} = \rho(t^*(\hat{\ell})).$$

Note also,

$$\hat{\ell} = \frac{f_1^*}{f_1^* + f_2^*} = \frac{1}{2 + (v_2/v_1)} < \frac{1}{3}.$$

That is, the robust weight puts more than twice as much emphasis on the Rawlsian (or egalitarian) objective as on the utilitarian objective. One can conclude that in general there is no “robustness equivalence principle,”

in the sense that substituting the robust parameter into the original scalarization (via maximization of the corresponding weighted objective) might *not* lead to a robust decision.²⁶

Example 9. Consider the robust allocation of two resources to $n = 2$ agents. The total amount of each resource has been normalized to one. Allocations are determined as decisions $x = (x_1, x_2) \in [0, 1]^2$, under which agent 1 obtains x_1 of resource 1 and x_2 of resource 2, whereas agent 2 obtains $1 - x_1$ of resource 1 and $1 - x_2$ of resource 2. Agent 1's utility is $f_1(x) = x_1^\alpha x_2^{1-\alpha}$, and agent 2's utility is $f_2(x) = (1 - x_1)^\beta (1 - x_2)^{1-\beta}$, where $\alpha, \beta \in (0, 1)$ are given scalars. Figure 5 provides an illustration in the corresponding Edgeworth box (see, e.g., Pareto 1906, p. 187). Because a robust allocation is necessarily Pareto-optimal, both agents' marginal rates of substitution for the two goods must be equal, so $x \in \mathcal{P}$ if and only if

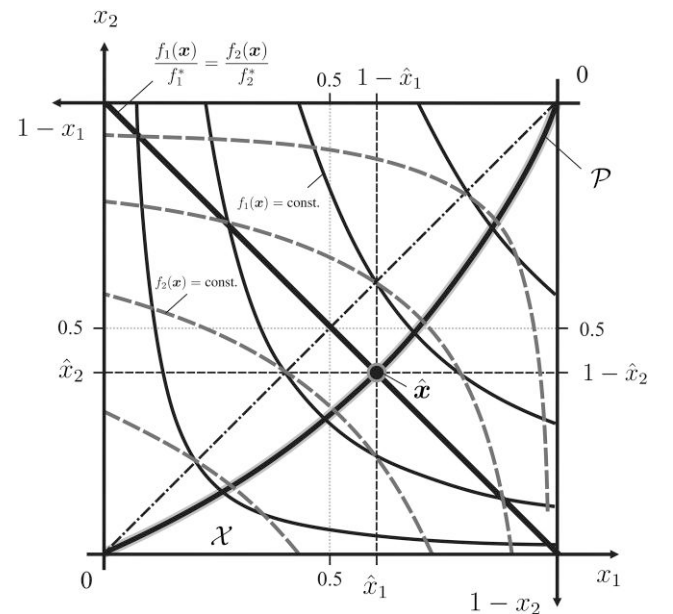
$$\frac{\partial f_1(x)/\partial x_1}{\partial f_1(x)/\partial x_2} = \frac{\alpha}{1 - \alpha} \frac{x_2}{x_1} = \frac{\beta}{1 - \beta} \frac{1 - x_2}{1 - x_1} = \frac{\partial f_2(x)/\partial x_1}{\partial f_2(x)/\partial x_2}. \quad (22)$$

In addition, a robust allocation must maximize the performance ratio, and one can verify that this requires the boundary performance ratios to be equal (i.e., a balancedness condition), so

$$f_1(x) = x_1^\alpha x_2^{1-\alpha} = (1 - x_1)^\beta (1 - x_2)^{1-\beta} = f_2(x), \quad (23)$$

because $f_1^* = f_2^* = 1$. In the interesting special case where

Figure 5. Robust Allocation of Resources $\hat{x} = (\hat{x}_1, \hat{x}_2)$ in Edgeworth Box $\mathcal{X} = [0, 1] \times [0, 1]$



Note. For two agents with utility functions (f_1 and f_2) specified in Example 9 (for $\alpha > \beta$), the unique robust allocation satisfies Pareto efficiency in Equation (22) and the balancedness condition in Equation (23), under full ambiguity.

$\alpha + \beta = 1$, Equations (22) and (23) together yield the unique robust solution $\hat{x} = (\alpha, \beta)$, resulting in an optimal performance index of $\rho^* = \alpha^\alpha \beta^\beta = \alpha^\alpha (1 - \alpha)^{1-\alpha}$, with the last expression being a strictly convex function in $\alpha \in (0, 1)$ achieving its minimum (of $1/2$) at $\alpha = 1/2$ and ρ^* going to one as $\alpha \rightarrow \{0^+, 1^-\}$.²⁷

3.7. Limited Ambiguity

Consider the (nonempty, compact) subset $\Omega \subset \mathbb{R}_+^n \setminus \{0\}$ which may reflect the decision maker's a priori knowledge about the relevant weights for the problem at hand, and let

$$\rho(x|\Omega) = \inf_{w \in \Omega} \varphi(x|w), \quad x \in \mathcal{X}, \quad (24)$$

be the corresponding Ω -conditioned performance index. The different weight combinations considered feasible may not necessarily be normalized, especially at an initial stage in practice when information about reasonable weight ranges is being compiled. The function $\pi: \mathbb{R}_+^n \setminus \{0\} \rightarrow \Delta$, with $\pi(w) = w/\|w\|_1$, describes the radial projection of any nonzero weight w in $\mathbb{R}_+^n$ onto the unit simplex Δ . It allows for the normalization under weight ambiguity; see Remark 2. The radial projection is a nonlinear function which is homogeneous of degree zero (i.e., $\pi(\alpha w) = \pi(w)$, for all $\alpha > 0$ and all $w \in \mathbb{R}_+^n \setminus \{0\}$). It can be represented using the conical closure (see, e.g., Berge 1963, p. 14), $\mathcal{C}(\Omega) = \{\hat{w} \in \mathbb{R}_+^n : \hat{w} = \alpha w, w \in \Omega, \alpha \in \mathbb{R}_+\}$, as described, among other useful properties of $\pi(\cdot)$, in the following auxiliary result.

Lemma 8. Let Ω be a (nonempty, compact) subset of $\mathbb{R}_+^n \setminus \{0\}$. Then (i) (a) $\pi(\Omega) = \mathcal{C}(\Omega) \cap \Delta$ is nonempty and compact, and (b) $\pi(\Omega) = \pi(\partial\Omega)$. (ii) If $\Omega' \subset \Omega$ is open (in $\mathbb{R}_+^n$), then $\pi(\Omega')$ is open (in Δ). (iii) If $\Omega' \subseteq \Omega$ and $\Omega' \neq \emptyset$, then $\partial\pi(\Omega') = \partial\pi(\partial\Omega')$ is nonempty and compact. (iv) If Ω is convex, then $\pi(\Omega)$ is convex. (v) If $\Omega' \subset \Omega$ is a straight line segment, then $\pi(\Omega')$ is a straight line segment (or a point). (vi) $\pi(\text{co}(\Omega)) = \text{co}(\pi(\partial\Omega))$. (vii) If $\Omega = \Omega' \cup \Omega''$ for some $\Omega', \Omega'' \subset \mathbb{R}_+^n$, then $\pi(\Omega) = \pi(\Omega') \cup \pi(\Omega'')$. (viii) If $\Omega = \Omega' \cap \Omega''$ for some $\Omega', \Omega'' \subset \mathbb{R}_+^n$, then $\pi(\Omega) = \mathcal{C}(\Omega' \cap \Omega'') \cap \Delta$.

Part (i) of Lemma 8 notes that (a) the radial projection of Ω onto Δ can be obtained by simply intersecting the conical closure of Ω with Δ , and (b) one can limit attention to the boundary $\partial\Omega$ (ignoring interior points of Ω). Although a continuous function generally does not map open sets to open sets (e.g., a constant function would not), part (ii) guarantees that $\pi(\cdot)$ does exactly that, provided the “interiority” of a point in Ω is assessed in $\mathbb{R}_+^n$ and then, after its projection, in the lower-dimensional Δ . Part (iii) ensures that weight ambiguity in *any* (nonempty) subset Ω' of Ω leads to a compact domain $\partial\pi(\partial\Omega')$ for a robust decision according to Proposition 5 below. Part (iv) establishes that convex sets are projected to convex sets.²⁸ Following (v) and (vi), the radial

projection leaves straight-line geometries intact, thus, for example, converting (bounded) polyhedra in $\mathbb{R}_+^n$ to polytopes in Δ . Finally, parts (vii) and (viii) note that the radial projection of a union of sets is the union of the corresponding single-set projections, but the same does generally not apply to an intersection.²⁹

Example 10. Assume the (nonnormalized) weights w considered reasonable by the decision maker are such that each of its components w_i , for $i \in \mathcal{N}$, is known to lie in some interval, resulting in a rectangular ambiguity set,

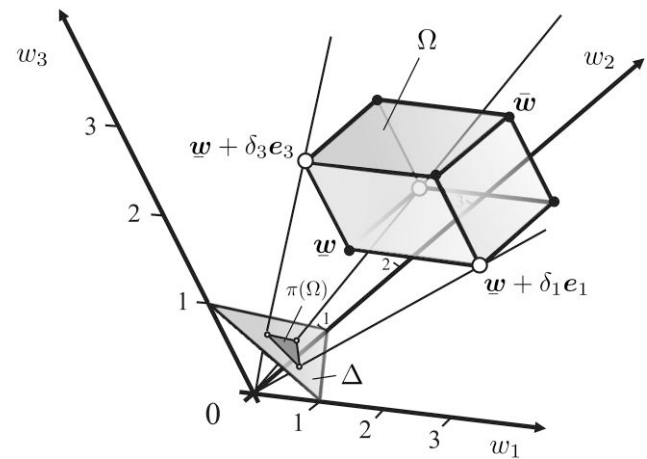
$$\Omega = [\underline{w}, \bar{w}] = \prod_{i \in \mathcal{N}} [\underline{w}_i, \bar{w}_i],$$

with bounds $\underline{w} = (\underline{w}_1, \dots, \underline{w}_n) \neq 0$ and $\bar{w} = (\bar{w}_1, \dots, \bar{w}_n) \geq \underline{w}$. Using Lemma 8, leveraging the preservation of straight-line geometries, it can be verified that the radial projection of Ω onto Δ has the form

$$\pi(\Omega) = \text{co} \left(\left\{ \frac{\underline{w} + \delta_1 e_1}{\|\underline{w}\|_1 + \delta_1}, \dots, \frac{\underline{w} + \delta_n e_n}{\|\underline{w}\|_1 + \delta_n} \right\} \right), \quad (25)$$

where $\delta = (\delta_1, \dots, \delta_n) = \bar{w} - \underline{w}$ and $\|\underline{w}\|_1 = \sum_{i=1}^n \underline{w}_i$. As an example, Figure 6 illustrates that the radial projection of a box in $n = 3$ dimensions is a triangle in Δ . Thus, the shape of $\pi(\Omega)$ depends on (at most) n of the original 2^n vertices of Ω . Moreover, if the modulus of the lower bound of the (nontrivial) box-shaped ambiguity set Ω becomes very small, then the situation tends to become equivalent to full ambiguity. To see this, consider a $\bar{w}$ with only positive components and let $\underline{w} = \varepsilon e_1$, for a sufficiently small $\varepsilon > 0$. Then, by Equation (25) it is $\pi(\Omega) = \text{co}(\{\lambda^{(1)}, \dots, \lambda^{(n)}\})$, where $\lambda^{(1)} = e_1$ and $\lambda^{(i)} = (\varepsilon e_1 + \bar{w}_i e_i)/(\varepsilon + \bar{w}_i)$, for $i \in \{2, \dots, n\}$. Hence, $\lim_{\varepsilon \rightarrow 0^+} \lambda^{(i)} = e_i$, for all $i \in \mathcal{N}$,

Figure 6. Radial Projection of Box-Shaped Ambiguity Set $\Omega = [\underline{w}, \bar{w}]$



Note. With $\underline{w} = (1, 1, 1)$ and $\bar{w} = (3, 2, 2)$, a radial projection onto Δ yields $\pi(\Omega) = \text{co}(\{(3/5, 1/5, 1/5), (1/4, 1/2, 1/4), (1/4, 1/4, 1/2)\})$, as described in Example 10.

implying full ambiguity in the limit, as $\Delta = \text{co}(\{e_1, \dots, e_n\}) = \pi(\mathbb{R}_+^n \setminus \{0\})$.

Under limited ambiguity, a coordinate-wise decomposition is no longer available. However, because $\varphi(x|\cdot)$ is quasiconcave, for any given $x \in \mathcal{X}$, as established in the proof of Proposition 2, the upper contour sets of the performance ratio in the space of weights are necessarily convex. This in turn allows restricting attention, for the representation of the performance index, to the extreme points of the convex hull of the ambiguity set, which yields a representation much in the same spirit as before.

Proposition 5. *The Ω -conditioned performance index in Equation (24) is such that*

$$\rho(x|\Omega) = \min_{\lambda \in \partial\pi(\partial\Omega)} \varphi(x|\lambda) = \min_{\lambda \in \Lambda} \varphi(x|\lambda), \quad x \in \mathcal{X},$$

for some “extremal base” Λ which is a (compact) subset of $\partial\pi(\partial\Omega)$, such that $\text{co}(\Lambda) = \text{co}(\pi(\Omega))$.

Although choosing $\Lambda = \partial\pi(\Omega) = \partial\pi(\partial\Omega)$ (cf. Lemma 8(i)(b)) is always possible, it is usually advantageous to opt for the smallest possible extremal base Λ , so it consists only of the “extreme points” of $\text{co}(\pi(\Omega))$, where the latter cannot be represented as convex combinations of other points in Λ (see, e.g., Rockafellar 1970, section 18, p. 162).³⁰ If Ω is finite or a (finite) polytope, then the smallest extremal base Λ of extreme points of $\text{co}(\pi(\Omega))$ is also finite.³¹ Because any bounded convex set can be approximated by a finite polytope (Bronstein 2008), a finite extremal base can be used to represent the convex hull of the given ambiguity set (or its radial projection onto Δ) up to any desired precision. That $\pi(\Omega)$ may be nonconvex is not important because the level sets of the performance ratio are convex, so that the minima of $\varphi(x|\cdot)$ on $\pi(\Omega)$ are attained on the boundary of its convexification, $\text{co}(\pi(\Omega))$.

Example 11. Consider the robust allocation of resources to $n = 2$ agents as in Example 9, given the box-shaped ambiguity set Ω as in Example 11, with its radial projection specified by Equation (25),

$$\pi(\Omega) = \text{co}(\Lambda) = \text{co}\left(\left\{\frac{(\bar{w}_1, \bar{w}_2)}{\bar{w}_1 + \bar{w}_2}, \frac{(\underline{w}_1, \underline{w}_2)}{\underline{w}_1 + \underline{w}_2}\right\}\right) \subsetneq \Delta,$$

where $\Lambda = \{\lambda^{(1)}, \lambda^{(2)}\}$. Thus, by virtue of Proposition 5 the Ω -conditioned performance index becomes

$$\rho(x|\Omega) = \min \left\{ \frac{\bar{w}_1 f_1(x) + \underline{w}_2 f_2(x)}{\max_{\hat{x} \in \mathcal{X}} \{\bar{w}_1 f_1(\hat{x}) + \underline{w}_2 f_2(\hat{x})\}}, \frac{\underline{w}_1 f_1(x) + \bar{w}_2 f_2(x)}{\max_{\hat{x} \in \mathcal{X}} \{\underline{w}_1 f_1(\hat{x}) + \bar{w}_2 f_2(\hat{x})\}} \right\}, \quad x \in \mathcal{X}.$$

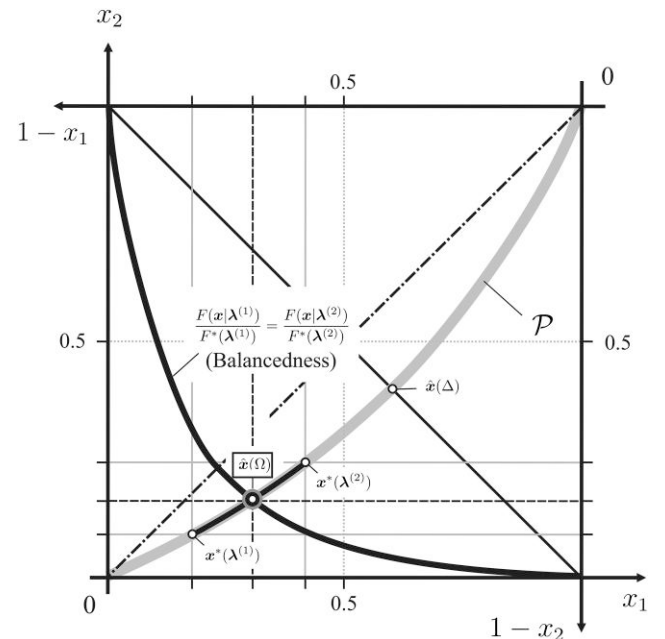
For a pseudo-robust decision x (in Ψ), balancedness must hold, so that, by substituting f_1, f_2 from Example 9,

$$\frac{F(x|\lambda^{(1)})}{F(x|\lambda^{(2)})} = \frac{\bar{w}_1 x_1^\alpha x_2^{1-\alpha} + \underline{w}_2 (1-x_1)^\beta (1-x_2)^{1-\beta}}{\underline{w}_1 x_1^\alpha x_2^{1-\alpha} + \bar{w}_2 (1-x_1)^\beta (1-x_2)^{1-\beta}} = \frac{F^*(\lambda^{(1)})}{F^*(\lambda^{(2)})}. \quad (26)$$

Meanwhile, a feasible allocation decision $x = (x_1, x_2)$ is Pareto-optimal (in $\mathcal{P}$) if and only if it satisfies Equation (22), as before. When the box-shaped ambiguity set becomes smaller, the elements of the extremal base Λ become more similar and coincide in the limit (i.e., $\lambda^{(1)} - \lambda^{(2)} \rightarrow 0$ as $\|\bar{w} - \underline{w}\|_1 \rightarrow 0^+$). Figure 7 shows the robust allocation $\hat{x} = \hat{x}(\Omega)$ on the contract curve between the Pareto-optimal allocations $x^*(\lambda^{(i)})$ that maximize the weighted objective $F(x|\lambda^{(i)})$ for $i \in \mathcal{N}$. Under full ambiguity (i.e., for $\Omega = \Delta$), it is $\lambda^{(i)} \equiv e_i$; see Example 9.

Example 12. In some applications, the weights are naturally ranked by importance (see, e.g., Wang and Fu 2020). The ordered set of weights, $\Omega' = \{(w_1, \dots, w_n) \in \mathbb{R}_+^n : w_1 \geq w_2 \geq \dots \geq w_n\} = \mathcal{C}(\Omega')$, has a radial projection onto Δ of the form $\pi(\Omega') = \text{co}(\{\lambda^{(i)} : i \in \mathcal{N}\})$, where $\lambda^{(i)} = \frac{1}{i} \sum_{j=1}^i e_j$, for all $i \in \mathcal{N}$. Combining this importance-ranking with a box-shaped ambiguity set $\Omega'' = [\underline{w}, \bar{w} + \delta]$ as in Example 10, with $\pi(\Omega'') = \text{co}(\{\lambda''^{(i)} : i \in \mathcal{N}\})$, where $\lambda''^{(i)} = (\underline{w} + \delta_i e_i) / (\|\underline{w}\|_1 + \delta_i)$,

Figure 7. Robust Allocation of Resources, $\{\hat{x}(\Omega)\} = \arg \max_{x \in \mathcal{P}} \rho(x|\Omega)$, to Two Agents in Example 11 (for $\alpha > \beta$)



Note. Under limited ambiguity $\Omega = [\underline{w}, \bar{w}]$ as in Example 10, the robust allocation $\hat{x}(\Omega)$ satisfies Pareto-efficiency in Equation (22) and balancedness in Equation (26); it differs from the allocation $\hat{x}(\Delta)$ under full ambiguity, which maximizes $\rho(x) = \rho(x|\Delta)$ on $\mathcal{P}$.

for all $i \in \mathcal{N}$, leads to

$$\Lambda = \left\{ \mathbb{1}_{\{\underline{w} + \delta_i e_i \notin \Omega'\}} \lambda^{(i)} + \mathbb{1}_{\{\underline{w} + \delta_i e_i \in \Omega'\}} \lambda''^{(i)} : i \in \mathcal{N} \right\}.$$

This extremal base (of cardinality n) suits practical applications where a decision maker disposes of plausible ranges for the weights to be placed on the criteria, together with a ranking of their importance (cf. Section 4.3.2).

Remark 12. (i) Limiting ambiguity can only increase robustness performance. That is, if $\Omega, \hat{\Omega} \subset \mathbb{R}_+^n \setminus \{0\}$ are nonempty and compact and satisfy $\pi(\hat{\Omega}) \subseteq \pi(\Omega)$, then $\rho(\cdot|\Omega) \leq \rho(\cdot|\hat{\Omega})$, which follows directly from the definition of the Ω -conditioned performance index in Equation (24). (ii) In the absence of ambiguity, the optimal robustness performance has to be maximal. That is, if $\hat{\Omega} = \{w\}$, then $\rho^*(\hat{\Omega}) = \max_{x \in \mathcal{X}} \rho(x|\{w\}) = F^*(w)/F^*(w) = 1$.

3.8. Criterion Ambiguity

Given a (nonempty, compact) action set $\mathcal{X} \subset \mathbb{R}^m$ as in Section 2, we now allow for ambiguity in each of the n' criteria, given by the continuous functions $g_i : \mathcal{X} \times \Theta \rightarrow \mathbb{R}_+$, for $i \in \mathcal{N}' = \{1, \dots, n'\}$, which map (x, θ) to real numbers, where $x \in \mathcal{X}$ is a feasible action and θ is an ex ante unknown state in the (nonempty, compact) state space $\Theta \subset \mathbb{R}^l$, and where $l, m, n' \geq 1$ are given integers. The unknown state θ represents the decision maker's uncertainty about the exact value that each of his n' objectives may attain under a chosen action x .

Remark 13. (i) Continuity of the criteria in the state implies that a small perturbation in the state can have only a small impact on the decision maker's objective. This continuity is automatically satisfied when the state space Θ is finite.³² (ii) The unknown state may introduce dependencies between the different g_i . Indeed, even if $\theta = (\theta_1, \dots, \theta_{n'})$ (for $l = n'$) and criteria are such that $g_i(x, \theta) \equiv g_i(x, \theta_i)$, for all $i \in \mathcal{N}'$, then the different criteria's ambiguity may still be linked via common constraints in Θ .³³ But if in addition the state space is a Cartesian product, so $\Theta = \Theta_1 \times \dots \times \Theta_{n'}$, and $\theta_i \in \Theta_i$, for all $i \in \mathcal{N}$, then the decoupling of ambiguity across the different criteria is complete, in the sense that observing the part of the state that determines one criterion does not reduce ambiguity for any other criterion.

To assess goal achievement of an action in a given state, the decision maker considers a scalarization of his multicriteria optimization problem by means of a weighted objective,

$$G(x, \theta|\lambda') = \sum_{i=1}^{n'} \lambda'_i g_i(x, \theta), \quad (x, \theta, \lambda') \in \mathcal{X} \times \Theta \times \Delta',$$

where $\Delta' = \{w' = (w_1, \dots, w_{n'}) \in \mathbb{R}_+^{n'} : w_1 + \dots + w_{n'} = 1\}$

denotes the set of all (normalized) weights. As in the nontriviality condition (N) (cf. Section 2), we assume that there exists a default decision (x_d) such that the decision maker's weighted objective is positive across all states, that is,

$$\exists x_d \in \mathcal{X} : G(x_d, \theta|\lambda') > 0, \quad (\theta, \lambda') \in \Theta \times \Delta'. \quad (\text{N}')$$

The modified nontriviality condition (N') ensures that the ex post optimal objective is always positive, so

$$G^*(\theta|\lambda') = \max_{x \in \mathcal{X}} G(x, \theta|\lambda') \geq G(x_d, \theta|\lambda') > 0,$$

for all $(\theta, \lambda') \in \Theta \times \Delta'$.

Remark 14. The modified nontriviality condition (N') can always be satisfied, without altering the set of ex post optimal decisions, by using a (positive, translated) criterion $\hat{g}_i = g_i + \varepsilon$, for some $\varepsilon > 0$; see also Remark 4.

Given a (nonempty, compact) ambiguity set $\Omega' \subset \mathbb{R}_+^{n'}$, the decision maker's robustness objective, as in Section 3.7, is to maximize the (Ω', Θ) -conditioned performance index,

$$\rho(x|\Omega', \Theta) = \min_{(\lambda', \theta) \in \pi(\Omega') \times \Theta} \frac{G(x, \theta|\lambda')}{G^*(\theta|\lambda')}, \quad x \in \mathcal{X}.$$

The following result recasts the robustness objective into a by-now-familiar representation.

Proposition 6. Let $\Lambda' \subset \Delta'$ be an extremal base of $\text{co}(\pi(\Omega'))$. Then

$$\rho(x|\Omega', \Theta) = \min_{\lambda' \in \Lambda'} \left\{ \min_{\theta \in \Theta} \frac{G(x, \theta|\lambda')}{G^*(\theta|\lambda')} \right\}, \quad x \in \mathcal{X}. \quad (27)$$

Based on Proposition 6, if the (smallest) extremal base is finite, we can reduce the general multicriteria decision problem to our basic framework in Section 2, with full weight ambiguity and no criterion ambiguity.

Corollary 1. Assume that there exists a (finite, smallest) integer $n \geq 1$ so that $\Lambda' = \{\lambda^{(1)}, \dots, \lambda^{(n)}\}$ and $\text{co}(\Lambda') = \text{co}(\pi(\Omega'))$, and let $f_i(x) = \min_{\theta \in \Theta} \{G(x, \theta|\lambda^{(i)})/G^*(\theta|\lambda^{(i)})\}$, for all $(x, i) \in \mathcal{X} \times \mathcal{N}$. Then

$$\rho(x) = \min_{i \in \mathcal{N}} \left\{ \frac{f_i(x)}{f_i^*} \right\}, \quad x \in \mathcal{X}, \quad (28)$$

where $f_i^* = \max_{x \in \mathcal{X}} f_i(x) = 1$, for all $i \in \mathcal{N} = \{1, \dots, n\}$, represents $\rho(\cdot|\Omega', \Theta)$ in Equation (27).

Example 13. (i) In the case of full ambiguity, it is $\Lambda' = \Delta'$, so that $n = n'$ and $\Delta' = \text{co}(\{e_1, \dots, e_{n'}\}) = \Delta$, with $f_i(x) = \min_{\theta \in \Theta} \{g_i(x, \theta)/G^*(\theta|e_i)\}$, for all $(x, i) \in \mathcal{X} \times \mathcal{N}$. (ii) Consider now a multicriteria decision problem under full ambiguity, with criteria of the form $g_i(x, \theta) = (\hat{g}_i(x))^{\theta_i}$, for all $i \in \mathcal{N}$, where the state $\theta = (\theta_1, \dots, \theta_n)$ is only known to lie in the (nonempty,

compact) set $\Theta \subset \mathbb{R}_+^n$ (with $\Theta \neq \{0\}$ to avoid trivialities), and where the functions $\hat{g}_i: \mathcal{X} \rightarrow \mathbb{R}_{++}$ are continuous. Let $\hat{g}_i^* = \max_{x \in \mathcal{X}} \hat{g}_i(x) > 0$ and $\bar{\theta}_i = \max_{\theta \in \Theta} \theta_i$; furthermore, set

$$f_i(x) = \left(\frac{\hat{g}_i(x)}{\hat{g}_i^*} \right)^{\bar{\theta}_i} = \min_{\theta \in \Theta} \left(\frac{\hat{g}_i^{\theta_i}(x)}{\max_{x \in \mathcal{X}} \hat{g}_i^{\theta_i}(x)} \right), \quad x \in \mathcal{X},$$

for all $i \in \mathcal{N}$, taking into account that ξ^θ decreases in θ , for all $\xi \in [0, 1]$. The extreme state $\bar{\theta} = (\bar{\theta}_1, \dots, \bar{\theta}_n)$ determines the robust choice, as it leads ceteris paribus to the smallest criterion values—thus following a precautionary principle. Then, by Corollary 1 we follow our basic framework, maximizing the performance index in Equation (28) with respect to all Pareto-optimal decisions in Equation (10). For instance, if $\Theta = \Delta$, the reduction to the weighted objective in Equation (1) obtains with $f_i = \hat{g}_i / \hat{g}_i^*$, for all $i \in \mathcal{N}$.

Remark 15. (i) The minimization over the state space in Example 13 reflects a robust, *precautionary* stance: the decision maker evaluates actions under the most adverse plausible realization of the state with respect to the achievable performance *ratio*, consistent with a notion of *relative* worst-case robustness. (ii) The treatment of criterion ambiguity in this section also bears a conceptual resemblance to models of DRO, where decisions are evaluated against worst-case distributions within a specified ambiguity set. In our framework, the uncertainty is not over probability distributions but over states $\theta \in \Theta$, with performance assessed under the worst-case realization of both the state and the weights. This can be viewed as a non-probabilistic analogue of DRO, where the ambiguity set Ω' over the weights plays a role similar to ambiguity sets over distributions in DRO models. In particular, the two-layer minimization in Equation (27)—over weights and states—mirrors the inner DRO minimization over distributions, highlighting a parallel structure between worst-case evaluation across distributional and multicriteria settings.

3.9. General Multicriteria Objectives

In certain practical applications, it may seem appealing to consider alternatives to the arithmetic mean in Equation (1), such as a harmonic or geometric mean, with suitable weights. One could even think of using a (weighted) power mean which would accommodate each of the earlier options as a special case; see Example 14 below, which illustrates *aggregation ambiguity*. Rather than commit to a particular functional form for aggregating multiple criteria, however, we propose a general class of multicriteria objectives that adhere to a set of reasonable axioms. These axioms reduce (in the case of uniform weights) to those known to characterize the quasiarithmetic mean. A key finding, notably, is that

maximizing this general weighted objective is entirely equivalent to maximizing the arithmetic mean objective in Equation (1), provided the criteria are suitably transformed.

Let $H: \mathcal{Y} \times \Delta \rightarrow \mathbb{R}$ denote a (continuous) multicriteria objective, where $\mathcal{Y} \subseteq \mathbb{R}^n$ is a (nonempty) domain such that $\hat{f}(\mathcal{X}) \subset \mathcal{Y}$, with the vector of criteria $\hat{f} = (\hat{f}_1, \dots, \hat{f}_n)$, and where $\Delta \subset \mathbb{R}_+^n$ is the unit simplex; see Section 2.1. For any decision $x \in \mathcal{X}$ and weight $\lambda \in \Delta$, the multicriteria objective produces an overall score $\hat{H}(x|\lambda) = H(\hat{f}(x)|\lambda)$, which the decision maker would like to maximize by choosing an appropriate decision—in the presence of (possibly limited) ambiguity about λ (and possibly also about $\hat{f}$) as discussed in Section 3.7 (and Section 3.8). To guide the form of a sufficiently flexible and interpretable objective, we posit five axioms (Axioms 1–5), which nest those proposed by Kolmogorov (1930) in his seminal work “On the notion of mean.”

Axiom 1 (Monotonicity). For all $y = (y_1, \dots, y_n) \in \mathcal{Y}$, $\hat{y} = (\hat{y}_1, \dots, \hat{y}_n) \in \mathcal{Y}$, $\lambda = (\lambda_1, \dots, \lambda_n) \in \Delta$, and $i \in \mathcal{N}$:

$$\begin{aligned} \hat{y} - y &= (\hat{y}_i - y_i)e_i, \hat{y}_i > y_i \\ \Rightarrow \quad &\begin{cases} H(\hat{y}|\lambda) > H(y|\lambda), & \text{if } \lambda_i > 0, \\ H(\hat{y}|\lambda) = H(y|\lambda), & \text{if } \lambda_i = 0. \end{cases} \end{aligned}$$

Axiom 2 (Symmetry). For all $y = (y_1, \dots, y_n) \in \mathcal{Y}$, $\lambda = (\lambda_1, \dots, \lambda_n) \in \Delta$, and $i, j \in \mathcal{N}$:

$$\begin{aligned} (\hat{y}, \hat{\lambda}) &= (y, \lambda) + (y_i - y_j, \lambda_i - \lambda_j)(e_j - e_i) \\ \Rightarrow \quad &H(\hat{y}|\hat{\lambda}) = H(y|\lambda). \end{aligned}$$

Axiom 3 (Reflexivity). $H(\alpha \mathbf{1}_n|\lambda) = \alpha$, for all $\alpha > 0$ (with $\alpha \mathbf{1}_n \in \mathcal{Y}$) and $\lambda \in \Delta$.³⁴

Axiom 4 (Associativity). For all $y = (y_1, \dots, y_n) \in \mathcal{Y}$, $\lambda = (\lambda_1, \dots, \lambda_n) \in \Delta$, and $m \in \mathcal{N} \setminus \{n\}$:³⁵

$$\begin{aligned} (\lambda_1, \dots, \lambda_m) \neq \mathbf{0}, \quad &H\left(\sum_{i=1}^m y_i e_i \middle| \frac{\sum_{i=1}^m \lambda_i e_i}{\sum_{i=1}^m \lambda_i}\right) = \alpha \\ \Rightarrow \quad &H(\alpha \mathbf{1}_m, y_{m+1}, \dots, y_n|\lambda) = H(y|\lambda). \end{aligned}$$

Axiom 5 (Coordinate Filter). $H(y|e_i) = y_i$, for all $y = (y_1, \dots, y_n) \in \mathcal{Y}$ and $i \in \mathcal{N}$.

The significance of these five basic requirements is as follows: *Monotonicity* (Axiom 1) means that as long as the weight component λ_i is positive, increasing the value y_i of criterion $i \in \mathcal{N}$ must also increase the weighted objective, whereas for $\lambda_i = 0$ the weighted objective becomes insensitive to criterion i ; *symmetry* (Axiom 2) requires that the weighted objective is invariant with respect to any joint permutation of indices belonging to criteria and their associated weights; *reflexivity* (Axiom 3) imposes score-consistency in the sense that if all criterion scores are identical, then that should also be the value of the weighted objective, no matter

what (normalized) weights are applied; in a similar vein, *associativity* (Axiom 4) postulates that replacing a group of inputs with their internal weighted average leaves the overall aggregation unchanged; finally, in order to guarantee that the weights provide a homotopic relation between all criteria in isolation, the weighted objective is a *coordinate filter* (Axiom 5) if it yields the i -th component of $\mathbf{y}$ when putting all weight on the i -th coordinate (i.e., for $\boldsymbol{\lambda} = \mathbf{e}_i$).

Based on Kolmogorov's result,³⁶ we define the *general h -mean*,

$$H(\mathbf{y}|\boldsymbol{\lambda}) = h^{-1}\left(\sum_{i=1}^n \lambda_i h(y_i)\right), \quad (29)$$

where the *kernel* $h : \mathcal{D} \rightarrow \mathbb{R}$ is a continuous strictly monotonic function, defined on a suitable domain $\mathcal{D} \subset \mathbb{R}$.

Lemma 9. The general h -mean $H : \mathcal{Y} \times \Delta \rightarrow \mathbb{R}$ in Equation (29) satisfies Axioms 1–5.

The following example illustrates the flexibility afforded by the general h -mean.

Example 14. Consider two well-known averages.

i. Let $\mathcal{Y} = \mathbb{R}_{++}^n$. For any weight $\boldsymbol{\lambda} \in \Delta$ and power parameter $p \in \mathbb{R}$, define the *power mean*:

$$M_p(\mathbf{y}|\boldsymbol{\lambda}) = \begin{cases} (\sum_{i=1}^n \lambda_i y_i^p)^{1/p}, & p \neq 0, \\ \prod_{i=1}^n y_i^{\lambda_i}, & p = 0, \end{cases}$$

which corresponds to Equation (29) with $h(y) = y^p$ for $p \neq 0$, and $h(y) = \log y$ for $p = 0$. This recovers the harmonic ($p = -1$), geometric ($p = 0$), and arithmetic ($p = 1$) means, among others. In the limit, the power mean approaches the minimum (for $p \rightarrow -\infty$) or maximum (for $p \rightarrow +\infty$) of all y_i with positive weights. Importantly, M_p is increasing in p (see, e.g., Hardy et al. 1934, theorem 16, p. 26) and uniquely satisfies homogeneity: $M_p(\alpha \mathbf{y}|\boldsymbol{\lambda}) = \alpha M_p(\mathbf{y}|\boldsymbol{\lambda})$ for all $\alpha > 0$ (see, e.g., Hardy et al. 1934, theorem 84, p. 68).³⁷

ii. For $h(\cdot) = \exp(\cdot)$, the general h -mean in Equation (29) becomes a weighted mean in the log-semiring, related to the LogSumExp (LSE) function, because $H(\mathbf{y}|\boldsymbol{\lambda}) = \text{LSE}(y_1 + \log \lambda_1, \dots, y_n + \log \lambda_n)$, for all $(\mathbf{y}, \boldsymbol{\lambda}) \in \mathbb{R}^n \times \text{int}(\Delta)$, where $\text{LSE}(\mathbf{y}) = \log(\sum_{i=1}^n \exp(y_i))$, for all $\mathbf{y} \in \mathbb{R}^n$. Such a formulation carries fruit in probabilistic modeling and neural networks, where LSE and softmax appear naturally because of their smoothness properties.³⁸

3.9.1. Reduction to Basic Framework. Crucially, the general h -mean in Equation (29) reduces to our standard weighted objective in Equation (1) under a transformation of the criteria, by setting $f_i = h \circ \hat{f}_i$, so that

$$\hat{H}(\mathbf{x}|\boldsymbol{\lambda}) = h^{-1}(F(\mathbf{x}|\boldsymbol{\lambda})), \quad (\mathbf{x}, \boldsymbol{\lambda}) \in \mathcal{X} \times \Delta.$$

If the kernel h is increasing (resp., decreasing), then maximizing $\hat{H}(\cdot|\boldsymbol{\lambda})$ is equivalent to maximizing (resp.,

minimizing) $F(\cdot|\boldsymbol{\lambda})$. Thus, our earlier results carry over directly—with minor adjustments in the decreasing case, as detailed in Section 3.9.2.

Example 15. Consider Example 13 (ii) for the “diagonal” state space $\Theta = \{\boldsymbol{\theta} \in \mathbb{R}_+^n : \theta_1 = \dots = \theta_n \in [\underline{p}, \bar{p}]\}$, where the constants $\underline{p}, \bar{p}$ are such that $0 < \underline{p} < \bar{p} < \infty$. By Example 14 (i), this represents a situation in which the decision maker is uncertain about the appropriate aggregation method, except that a (homogeneous) power mean M_p should be used, for some $p \in [\underline{p}, \bar{p}]$. The result in Example 13 (ii) suggests as robust choice the largest available option: the power mean $M_{\bar{p}}$.

Remark 16. (i) As Example 15 suggests, the insights about criterion ambiguity in Section 3.8 may sometimes be combined with the general representation of multicriteria objectives in Section 3.9 to handle *aggregation ambiguity*, that is, uncertainty over which aggregation rule or kernel should be used. (ii) For the general h -mean in Equation (29), we define $H = h^{-1} \circ F_h$ with $F_h = \sum_{i=1}^n \lambda_i (h \circ f_i)$. However, in general, $h^{-1} \circ (F_h/F_h^*) \neq (h^{-1} \circ F_h)/(h^{-1} \circ F_h^*)$ with $F_h^*(\cdot) = \max F_h(\mathcal{X}|\cdot)$, so our optimization focuses on F_h directly. (iii) For the practically very important power mean in Example 14 (i) and Example 15 (as long as $p \neq 0$), the *robust* optimization of F_h is equivalent to the *robust* optimization of M_p , because $h(\cdot) = (\cdot)^p$ and $h^{-1}(\cdot) = (\cdot)^{1/p}$ both feature multiplicative separability.

3.9.2. Special Case: Minimization Under Decreasing Kernel.

For a decreasing kernel, $f_i = h \circ \hat{f}_i$ transforms larger $\hat{f}_i$ into smaller f_i , so that the objective reflects a *smaller-is-better* interpretation (akin to optimizing a loss function). Assuming $\mathbb{R}_{++}^n \subset \mathcal{Y}$, by Axiom 3 it is $\lim_{\alpha \rightarrow 0^+} H(\alpha \mathbf{1}_n|\boldsymbol{\lambda}) = H(\mathbf{0}^+|\boldsymbol{\lambda}) = 0$. Because in most practical applications (e.g., for logarithmic or inverse transformations in multiobjective loss minimization) it is $h(\mathbf{0}^+) = \infty$, we obtain $h^{-1}(\infty) = \mathbf{0}^+$, so that on the compact action set $\mathcal{X}$ it is $f_i^* = \min f_i(\mathcal{X}) = h(\hat{f}_i^*) > 0$, for all $i \in \mathcal{N}$, as $\hat{f}_i^* = \max \hat{f}_i(\mathcal{X})$ is necessarily finite. Hence, we can consider the *adjusted performance ratio*,

$$\varphi(\mathbf{x}|\boldsymbol{\lambda}) = \frac{F^*(\boldsymbol{\lambda})}{F(\mathbf{x}|\boldsymbol{\lambda})} \in (0, 1], \quad (\mathbf{x}, \boldsymbol{\lambda}) \in \mathcal{X} \times \Delta, \quad (30)$$

where $F^*(\boldsymbol{\lambda}) = \min F(\mathcal{X}|\boldsymbol{\lambda})$, for all $\boldsymbol{\lambda} \in \Delta$. A representation of the performance index $\rho(\mathbf{x})$ in Equation (4) then obtains, analogous to Equation (6), for all $\mathbf{x} \in \mathcal{X}$, as the minimum of the (adjusted) criterion-specific performance ratios $\phi_i(\mathbf{x}) = f_i^*/f_i(\mathbf{x})$ with respect to $i \in \mathcal{N}$. This defines the (nonempty, compact) set of pseudo-robust actions,

$$\Psi = \arg \max_{\mathbf{x} \in \mathcal{X}} \left\{ \min_{i \in \mathcal{N}} \frac{f_i^*}{f_i(\mathbf{x})} \right\}.$$

The set of Pareto-optimal actions $\mathcal{P}$ needs to be based on the original vector of criteria $\hat{\mathbf{f}}$, so that

$$\mathcal{P} = \{\mathbf{x} \in \mathcal{X} : (\hat{\mathbf{f}}(\mathbf{x}) \leq \hat{\mathbf{f}}(\mathbf{x}') \Rightarrow \hat{\mathbf{f}}(\mathbf{x}) = \hat{\mathbf{f}}(\mathbf{x}')), \mathbf{x}' \in \mathcal{X}\},$$

analogous to Equation (10). The robust decision set is then $\mathcal{R} = \Psi \cap \mathcal{P}$ as before, in Equation (11).

3.10. Applying the Method

“What do I do? What do I get? How do I adapt it to my case?”—We now address these practitioner questions by outlining the method’s core, the interpretation of its results, and important extensions for customization, as a navigation device for approaching the multicriteria decision problem introduced in Section 2.1 and its generalizations.

I. (Core) To determine an approximately robust decision $\hat{x}_\varepsilon$ (at any prespecified precision ε) with respect to n criteria, one needs to solve just $n+1$ optimization problems: one to maximize the ε -augmented performance index Φ_ε in Equations (12) and (13) and n to compute the normalization constants f_i^* in Equation (5) for the criterion-specific performance ratios ϕ_i in Equation (6). The associated normalized robust weight $\hat{\lambda} \in \Delta$, which encapsulates the tradeoffs between the criteria embedded in both the action set and the shape of the criterion functions (independent of any scaling), is then obtained (approximately) from Equations (19) and (20) by setting $\hat{x} = \hat{x}_\varepsilon$.

II. (Interpretation) The performance index $\rho(\hat{x}_\varepsilon) \in [0, 1]$, with $\rho(\cdot)$ given in Equation (6), guarantees a minimum percentage that the robust solution $\hat{x}_\varepsilon$ achieves of the optimal weighted objective in Equation (1), for any weight in Δ . This means that the performance of $\hat{x}_\varepsilon$ is guaranteed to be no worse than $\rho(\hat{x}_\varepsilon)$ times the best achievable performance under any possible weight vector, ensuring robustness to unknown or contested preferences. In this manner, the weight-induced subjectivity in multicriteria optimization can be removed. The robust weight rationalizes the robust decision as an optimum of the weighted objective in Equation (1) for $\lambda = \hat{\lambda}$. The vector $\hat{\lambda}$ can be interpreted as the revealed weight structure that best justifies the robust solution, based on the problem’s internal tradeoffs.

III. (Customization) The method accommodates several practically important extensions:

– *Partial weight information:* Prior knowledge or constraints on weights can be imposed by requiring weight vectors (not necessarily normalized) to belong to a suitable subset (cf. Section 3.7).

– *State-dependent criteria:* Uncertain or context-dependent criteria f_i can be captured by worst-case performance ratios, effectively reducing the problem to the core framework (cf. Section 3.8).

– *Generalized aggregation:* Rather than relying on the arithmetic mean in Equation (1), one may adopt alternative scalarization functions consistent with an axiomatic foundation (cf. Section 3.9).

Overall, the relatively robust methodology provides a computationally tractable and conceptually transparent

toolkit for balancing multiple, possibly ambiguous objectives in diverse real-world settings (cf. Section 4.3).

4. Discrete Applications

4.1. Finite Action Set

4.1.1. Relative Robustness Criterion. Consider a finite set of alternatives $\mathcal{X} = \{1, \dots, J\}$, evaluated according to n positive criteria $f_1, \dots, f_n : \mathcal{X} \rightarrow \mathbb{R}_{++}$. Each alternative $j \in \mathcal{X}$ receives a score $f_i(j) = s_{ij} > 0$, for all $i \in \mathcal{N}$. If we set $\hat{s}_i = \max_{j \in \mathcal{X}} s_{ij}$, then the performance index for the j -th option becomes

$$\rho_j = \min_{i \in \mathcal{N}} \frac{s_{ij}}{\hat{s}_i}, \quad j \in \mathcal{X}. \quad (31)$$

By Proposition 3, a robust decision $j^* \in \mathcal{R}$ can be found by computing the lower limit, as follows:

$$j^* \in \underline{\text{Lim}}_{\varepsilon \rightarrow 0^+} \arg \max_{j \in \mathcal{X}} \left\{ (1 - \varepsilon) \rho_j + \frac{\varepsilon}{n} \sum_{i \in \mathcal{N}} \frac{s_{ij}}{\hat{s}_i} \right\}. \quad (32)$$

This robust decision achieves the optimal performance index,

$$\rho^* = \max_{j \in \mathcal{X}} \rho_j = \rho_{j^*} = \min_{i \in \mathcal{N}} \frac{s_{ij^*}}{\hat{s}_i}. \quad (33)$$

As discussed in Section 3, simply maximizing the performance index ρ_j yields pseudo-robust solutions, which are generically inefficient; see, for instance, Example 6. We now compare the proposed approach with several alternative robustness criteria.

4.1.2. Alternative Robustness Criteria. The following three robustness criteria, defined in the present context of discrete action sets, are frequently used in the literature for dealing with parameter ambiguity.

i. The *Laplace criterion* corresponds to a weighted objective $F(\cdot | \lambda_{\text{Laplace}})$ in Equation (1) with a uniform weight, $\lambda_{\text{Laplace}} = (1, \dots, 1)/n \in \Delta$, leading to the “Laplace solution,”

$$j_{\text{Laplace}} \in \arg \max_{j \in \mathcal{X}} F(j | \lambda_{\text{Laplace}}) = \arg \max_{j \in \mathcal{X}} \sum_{i=1}^n s_{ij}. \quad (34)$$

ii. The *worst-case criterion* is defined as the minimum payoff across the admissible weights, which yields a so-called “maximin solution,”

$$j_{\text{WC}} \in \arg \max_{j \in \mathcal{X}} \min_{\lambda \in \Delta} F(j | \lambda) = \arg \max_{j \in \mathcal{X}} \min_{i \in \mathcal{N}} s_{ij}. \quad (35)$$

iii. The *absolute-regret criterion* evaluates absolute regret, that is, the maximum difference ex post between what is and the best that could have been, resulting in the “(minimax) absolute-regret solution,”

$$\begin{aligned} j_{\text{AR}} &\in \arg \min_{j \in \mathcal{X}} \max_{\lambda \in \Delta} \{F^*(\lambda) - F(j | \lambda)\} \\ &= \arg \min_{j \in \mathcal{X}} \max_{i \in \mathcal{N}} \{\hat{s}_i - s_{ij}\}. \end{aligned} \quad (36)$$

4.1.3. Comparison. The following example, which features a finite action set, illustrates the proposed robust

solution in Equation (32) against decisions recommended by alternative robustness criteria, notably the Laplace criterion in Equation (34), the WC (or maximin) solution in Equation (35), and the solution minimizing the (maximum) AR in Equation (36).

Example 16. Consider a discrete-choice situation for $J = 5$ options, for which the scores s_{ij} across $n = 3$ criteria are recorded in Table 1. In the context of relative robustness, options 2 and 5 are tied for the highest performance index in Equation (31) and are thus both pseudo-robust, so $\Psi = \{2, 5\}$. At the same time, option 5 Pareto-dominates option 2 (with $\mathcal{P} = \{1, 3, 4, 5\}$), so that $j^* = 5 \in \mathcal{R} = \Psi \cap \mathcal{P}$ is the unique robust choice. This solution can also be obtained directly from Equation (32), written in the form

$$j^* \in \underline{\text{Lim}}_{\varepsilon \rightarrow 0^+} \arg \max_{j \in \mathcal{X}} \Phi_\varepsilon(j),$$

where $\Phi_\varepsilon(j) = (1 - \varepsilon)\rho_j + (\varepsilon/n) \sum_{i=1}^n (s_{ij}/\hat{s}_i)$, for all $j \in \mathcal{X}$ and all $\varepsilon \in [0, 1]$. Indeed, because $\rho_2 = \rho_5 = 3/11 > \rho_1 = \rho_4 = 4/15 > \rho_3 = 1/24$, for sufficiently small $\varepsilon \in (0, 1]$ the ε -augmented performance index is

$$\max_{j \in \{2, 5\}} \Phi_\varepsilon(j) > \Phi_\varepsilon(4) > \Phi_\varepsilon(3).$$

Meanwhile, it is

$$\Phi_\varepsilon(5) - \Phi_\varepsilon(2) = \frac{\varepsilon}{3} \left(\frac{12-7}{24} + \frac{9-8}{15} \right) = \frac{\varepsilon}{3} \left(\frac{11}{40} \right) > 0, \quad \varepsilon \in (0, 1].$$

Therefore, $\mathcal{R}_\varepsilon = \{5\}$, for all sufficiently small $\varepsilon > 0$, so $j^* \in \underline{\text{Lim}}_{\varepsilon \rightarrow 0^+} \mathcal{R}_\varepsilon = \{5\}$, and one finds $j^* = 5$ as the unique robust option. Regarding alternative robustness evaluation, both the Laplace criterion in Equation (34) and the absolute-regret criterion in Equation (36) produce the solution $j_{\text{Laplace}} = j_{\text{AR}} = 3$ with very poor performance in at least one criterion. In fact, changing the payoff s_{13} from one to zero would produce even a zero performance index and zero worst-case performance guarantee for that same decision (still optimal under these criteria),³⁹ which may be rather difficult to justify in any real-world scenario. Finally, maximizing the worst-criterion performance in Equation (35) leads to indifference between options 1 and 4, at a suboptimal performance index and largest absolute regret. By contrast, the proposed robust solution $j^* = 5$ provides (by construction)

the best performance index and a reasonable compromise solution in terms of the other criteria. It guarantees a strictly positive performance across all criteria. From Equations (19) and (20) we find that the corresponding (normalized) robust weight vector is

$$\hat{\lambda} = \left(\frac{1}{12} + \frac{1}{6} + \frac{1}{36} \right)^{-1} \left(\frac{1}{12}, \frac{1}{6}, \frac{1}{36} \right) = (0.3, 0.6, 0.1).$$

Consistent with Equation (21), it is such that

$$\begin{aligned} \hat{\lambda}_{i s_{ij^*}} &= \left(\sum_{i \in \mathcal{N}} \frac{1}{s_{ij^*}} \right)^{-1} = \left(\frac{1}{12} + \frac{1}{6} + \frac{1}{36} \right)^{-1} = 3.6 \\ &\geq \max_{j \in \mathcal{X} \setminus \{j^*\}} \min_{i \in \mathcal{N}} \hat{\lambda}_i s_{ij}, \quad i \in \{1, 2, 3\}, \end{aligned}$$

where 3.6 corresponds to the harmonic mean of the criterion scores for the robust option j^* . Consider now the robustly reweighted scores $\sigma_{ij} = \hat{\lambda}_i s_{ij}$, for all $(i, j) \in \mathcal{N} \times \mathcal{X}$, displayed in Table 2, together with the associated robustness evaluations. It is interesting that after reweighting, option 3 goes from lowest to largest absolute regret, which highlights the sensitivity of the AR criterion to uncertainty about the relative importance of the criteria. The maximin payoff (WC) solution shifts from options 1 and 4 to option 5, whereas the Laplace solution shifts from option 3 to option 4. By contrast, the performance index remains unaffected by the robust reweighting (or any other reweighting), so option 5 is still the unique robust solution.

4.2. Data-Driven Approach

Consider any (nonempty, compact) action set $\mathcal{X} \subset \mathbb{R}^m$ for some integer $m \geq 1$, as in Section 2. Assume further that the decision maker does not know the functional form of the multicriterion $f = (f_1, \dots, f_n)$, but is still able to observe its value $s^k = (s_1^k, \dots, s_n^k)$ for different decisions $x^k \in \mathcal{X}$ over the course of $K \geq 1$ experiments, where $k \in \mathcal{K} = \{1, \dots, K\}$. Indeed, given the realized score set $\hat{\mathcal{S}} = \{s^k : k \in \mathcal{K}\}$ and the realized action set $\hat{\mathcal{X}} = \{x^k : k \in \mathcal{K}\} \subset \mathcal{X}$, the average criterion scores are⁴⁰

$$\hat{f}_i(x) = \frac{\sum_{k \in \mathcal{K}} \mathbf{1}_{\{x^k=x\}} s_i^k}{\sum_{k \in \mathcal{K}} \mathbf{1}_{\{x^k=x\}}}, \quad (x, i) \in \hat{\mathcal{X}} \times \mathcal{N},$$

for all $(x, i) \in \hat{\mathcal{X}} \times \mathcal{N}$. With this, the (data-driven) criterion-

Table 1. Discrete Decision Options in Example 16, Evaluated with Different Robustness Criteria

Option (j)	Criteria			Performance ratios/index				Alternative evaluations		
	s_{1j}	s_{2j}	s_{3j}	$s_{1j}/\hat{s}_1$	$s_{2j}/\hat{s}_2$	$s_{3j}/\hat{s}_3$	ρ_j	Laplace	WC	AR
1	24	20	16	1	10/11	4/15	4/15	20	16	44
2	7	6	32	7/24	3/11	8/15	3/11	15	6	28
3	1	2	60	1/24	1/11	1	1/24	21	1	23
4	22	22	16	11/12	1	4/15	4/15	20	16	44
5	12	6	36	1/2	3/11	3/5	3/11	18	6	24

Note. Values in **bold** indicate optimality for the corresponding robustness criterion.

Table 2. Discrete Decision Options in Example 16, Evaluated after Robust Reweighting

Option (j)	Criteria			Performance ratios/index				Alternative evaluations		
	σ_{1j}	σ_{2j}	σ_{3j}	$\sigma_{1j}/\hat{\sigma}_1$	$\sigma_{2j}/\hat{\sigma}_2$	$\sigma_{3j}/\hat{\sigma}_3$	$\hat{\rho}_j$	Laplace	WC	AR
1	7.2	12	1.6	1	10/11	4/15	4/15	20.8	1.6	4.4
2	2.1	3.6	3.2	7/24	3/11	8/15	3/11	8.9	2.1	9.6
3	0.3	1.2	6	1/24	1/11	1	1/24	7.5	0.3	12
4	6.6	13.2	1.6	11/12	1	4/15	4/15	21.4	1.6	4.4
5	3.6	3.6	3.6	1/2	3/11	3/5	3/11	10.8	3.6	9.6

Note. Values in **bold** indicate optimality for the corresponding robustness criterion.

specific performance ratios are given by

$$\hat{\phi}_i(x) = \frac{\hat{f}_i(x)}{\hat{f}_i^*}, \quad (x, i) \in \hat{\mathcal{X}} \times \mathcal{N},$$

where $\hat{f}_i^* = \max_{x \in \hat{\mathcal{X}}} \hat{f}_i(x)$ is the i -th maximum average criterion score. The (data-driven) performance index,

$$\hat{\rho}(x) = \min_{i \in \mathcal{N}} \hat{\phi}_i(x), \quad x \in \hat{\mathcal{X}},$$

is defined for all observed decisions. A (data-driven) robust decision $\hat{x}^*$ (among all sampled actions in $\hat{\mathcal{X}}$) can then be determined in the standard way by applying Proposition 3 to the observational analogues of our standard approach:

$$\hat{x}^* \in \underline{\text{Lim}}_{\varepsilon \rightarrow 0^+} \arg \max_{x \in \hat{\mathcal{X}}} \left\{ (1 - \varepsilon) \hat{\rho}(x) + \frac{\varepsilon}{n} \sum_{i \in \mathcal{N}} \hat{\phi}_i(x) \right\}. \quad (37)$$

In other words, $\hat{\rho}(x)$ captures the worst-case relative performance of action x across all observed criteria, and $\hat{x}^*$ identifies a robust decision, which is efficient and has the best-possible worst-case performance. Note that instead of going through the motions of actually taking the limit, it is sufficient to maximize the ε -augmented performance index, that is, the maximand in Equation (37), for a sufficiently small $\varepsilon \in (0, 1]$. This approximates the optimal robustness performance, captured by the optimal performance index ρ^* , arbitrarily closely by virtue of the ε -performance guarantee provided in Proposition 4. Finally, as in Equations (19) and (20), the (data-driven) robust weight is

$$\hat{\lambda} = \left(\sum_{i \in \mathcal{N}} \frac{1}{\hat{f}_i(\hat{x}^*)} \right)^{-1} \left(\frac{1}{\hat{f}_i(\hat{x}^*)} \right)_{i=1}^n \in \Delta. \quad (38)$$

Based on the observed data, this vector provides a robust estimate of the tradeoffs between the different criteria from the vantage point of relative robustness.

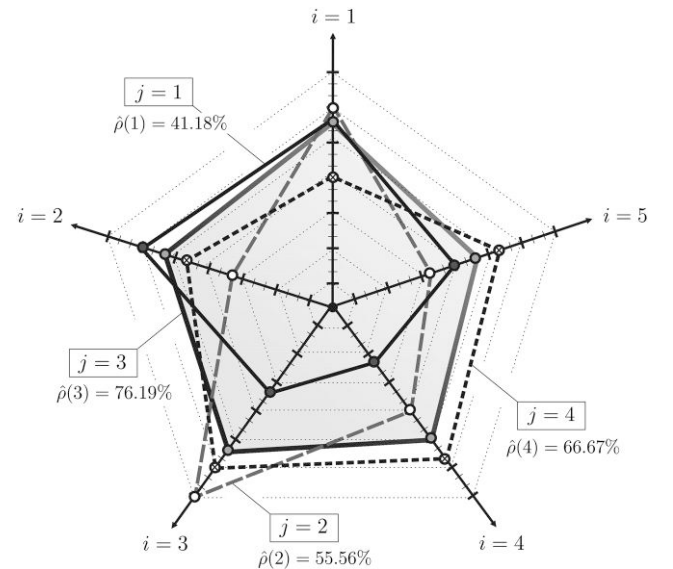
Example 17. Consider the joint evaluation of human performance on a certain task (such as giving a university lecture) by K evaluators who score J individuals in N different performance dimensions (or criteria) on a Likert scale (from one to seven),⁴¹ so $\mathcal{K} = \{1, \dots, K\}$, $\mathcal{N} = \{1, \dots, N\}$, and $\mathcal{X} = \{1, \dots, J\}$. Figure 8 shows a spider plot comparing $J = 4$ different individuals. Assuming that all evaluators scored all the individuals using a

seven-point Likert scale, the realized score set $\hat{\mathcal{S}}$ has elements $s^k \in \{1, \dots, 7\}^N$, for all $k \in \mathcal{K}$, whereas the realized action set $\hat{\mathcal{X}}$ is equal to $\mathcal{X}$. Individual 3, although never achieving a highest criterion-specific average score, exhibits the best data-driven performance index ($\hat{\rho}^* = \hat{\rho}(3) = 76.19\% = \max_{j \in \{1, 2, 3, 4\}} \hat{\rho}(j)$). Because in a noisy observation environment ties in the performance index are fairly unlikely, the set of pseudo-robust alternatives Ψ tends to be a singleton, thus also resulting in a singleton set of robust options (e.g., $\mathcal{R} = \{3\}$). This example illustrates how the data-driven approach identifies robust performance through balanced tradeoffs rather than peak performance in any one dimension, thus favoring individual 3, whose weakest dimension is relatively strong.

4.3. Real-World Applications

Relatively robust multicriteria optimization provides a flexible and transparent framework for decision-making under uncertainty. As developed in Sections 2

Figure 8. Average Scores for $N = 5$ Criteria (Evaluated by $K = 3$ Experts), and Resulting Data-Driven Performance Indices for $J = 4$ Options in Example 17



Note. All options are Pareto-efficient (i.e., $\mathcal{P} = \{1, 2, 3, 4\}$); only option 3 is also pseudo-robust (i.e., $\Psi = \{3\}$), so $\mathcal{R} = \Psi \cap \mathcal{P} = \{3\}$.

and 3, it is especially well suited for complex problems where the relative importance of multiple criteria is ambiguous or hard to justify. The method systematically balances competing objectives using only minimal assumptions on the decision maker's preferences, and delivers solutions that are both Pareto-efficient and robust (cf. Proposition 3). We illustrate its practical relevance in three diverse contexts: the energy trilemma, quality-adjusted life year (QALY)-based health evaluations, and corporate resource allocation. In each case, the relatively robust methodology facilitates structured tradeoff management and data-informed robustness, whether operating in an abstract policy space or on a finite empirical action set (cf. Sections 4.1 and 4.2).

4.3.1. Energy Trilemma: Balancing Energy Security, Equity, and Sustainability. The World Energy Council's Energy Trilemma framework evaluates nations on three key dimensions, namely "energy security" (i.e., the reliability and resilience of energy supply), "energy equity" (i.e., the accessibility and affordability of energy for all citizens), and "environmental sustainability" (i.e., the reduction of greenhouse gas emissions and environmental impact). Nations face challenges in improving one dimension without compromising others (World Energy Council 2024).⁴² For instance, increasing energy equity by subsidizing fossil fuels may undermine sustainability, while prioritizing environmental sustainability through renewable energy investment might initially reduce energy security or equity. An optimal strategy depends on the relative importance attributed to each of the three criteria. In this setting, a relatively robust approach can

- **Identify Robust Policies:** By modeling each nation's energy policies and outcomes as feasible decisions, the method identifies those policies that achieve strong tradeoffs across all three dimensions. For example, a relatively robust policy might balance investment in renewable energy, grid modernization, and subsidies for low-income households, ensuring consistent performance regardless of variations of the relative importance across dimensions, which might still be importance-ranked (cf. Section 3.7).

- **Quantify Tradeoffs:** The approach provides a clear representation of tradeoffs, such as how much equity might need to be sacrificed for a given improvement in sustainability under different weighting scenarios.

- **"Robustify" Strategic Goals:** Policymakers can develop long-term strategies that are resilient to evolving societal priorities, such as shifts toward greater emphasis on the precautionary principle (cf. Section 3.9).

4.3.2. QALY Impact of Diseases: Evaluating Quality-Adjusted Life Years. QALYs measure the impact of diseases and medical treatments by combining longevity and quality of life into a single index (Pliskin 1974, Pliskin

and Beck 1976, Miyamoto et al. 1998).⁴³ The impact of different impairments—such as mobility loss, chronic pain, or cognitive decline—depends on how these are weighted relative to each other in calculating overall health outcomes. Subject to restrictions on weights (e.g., a priority ranking; cf. Example 12) and with a flexible aggregation of criteria (cf. Section 3.9) the method can

- **Handle Uncertain Weightings:** When precise utility-weightings for different impairments are unavailable (as one would generally assume), the method identifies health interventions or treatment plans that remain effective across a wide range of subjective assessments. For example, it might suggest treatments that optimize outcomes for both pain relief and mobility restoration, ensuring robust quality-of-life improvements.

- **Optimize Resource Allocation:** Health authorities can use the method to prioritize interventions that deliver the highest QALY improvements per dollar spent, even when the relative importance of various health dimensions (e.g., physical versus mental health) is uncertain.⁴⁴

- **Support Evidence-Based Policy:** The method provides a performance guarantee for proposed policies, demonstrating their effectiveness regardless of the relative emphasis placed on specific impairments.

4.3.3. Resource Allocation in Companies: Balancing Risk, Resources, and Rival Assets. In organizations, resource allocation involves deciding how to distribute limited resources—financial, human, or physical—across competing projects. Each project is evaluated based on multiple criteria, such as "risk of noncompletion" (i.e., the probability that a project fails because of delays or unforeseen issues), "use of human resources" (i.e., the availability and workload of skilled personnel required for the project), and "use of rival assets" (i.e., the occupation of assets that cannot be used by multiple projects simultaneously, for example, specialized machinery). Applying the framework of relatively robust decision making allows organizations to⁴⁵

- **Identify Balanced Portfolios:** By modeling projects and their criteria as feasible decisions $x \in \mathcal{X}$, with performance criteria $f_1(x), \dots, f_n(x)$, the method selects a portfolio of projects that balances competing objectives, ensuring efficient resource use even when preferences over criteria are unclear.

- **Manage Tradeoffs:** By aggregating criterion-specific performance ratios $\phi_i(x)$ into a performance index $\rho(x)$, the method reveals, for example, how prioritizing low-risk projects impacts resource utilization and asset deployment, informing decisions about a robust balance between risk and resource efficiency.

- **Adapt to Uncertainty:** As organizational priorities shift (e.g., toward innovation or risk aversion), the method ensures that allocation strategies can remain robust across changing criteria (cf. Section 3.8) and

weight configurations for the criteria, and possibly across different aggregation methods (cf. Section 3.9).

In settings with discrete project sets and empirical observations, a data-driven approach can be applied (cf. Section 4.2).

5. Conclusion

Multicriteria optimization aims to resolve conflicts between competing objectives by finding Pareto-efficient decisions for which improving one criterion necessarily degrades at least one other. Although Pareto-efficiency sets a minimum standard of nonwastefulness, practical decision making often requires selecting a single compromise solution from the available Pareto frontier. Scalarization, such as assigning weights to criteria, is a common method for operationalizing this selection, but it assumes knowledge of the weights, which may not be available or easily justifiable, and it may also exclude Pareto-efficient solutions when the action set is nonconvex.

Here, we seek to mitigate the dependence on precise weight specifications by introducing a performance index that evaluates the worst-case weighted performance of a decision relative to its maximum potential. A Pareto-efficient decision that maximizes this index is viewed as relatively robust, balancing competing criteria in a way that offers resilience to weight uncertainty. A critical feature of the method is its computational simplicity, relying only on the compactness of the feasible set and the continuity of criterion functions to guarantee the existence and basic regularity of the solutions. This ensures broad applicability, even for complex, nonconvex problems. Criterion ambiguity and more general aggregation methods can be accommodated. In the case of finite sets, this method is completely general and can be operationalized through a data-driven approach under virtually no assumptions.

Appendix. Proofs

Proof of Proposition 1. The proof proceeds by analyzing what happens when either one of the two inequalities is violated, and subsequently when both are violated. This yields three cases, each of which leads to a contradiction.

- **Case 1:** If $f_i(\hat{x}) < f_i(x)$ and $f_j(\hat{x}) \leq f_j(x)$, then

$$F^*(\hat{\lambda}) = F(\hat{x}|\hat{\lambda}) = \sum_{i=1}^n \hat{\lambda}_i f_i(\hat{x}) < \sum_{i=1}^n \hat{\lambda}_i f_i(x) = F(x|\hat{\lambda}),$$

which is a contradiction.

- **Case 2:** If $f_i(\hat{x}) \geq f_i(x)$ and $f_j(\hat{x}) > f_j(x)$, then

$$F^*(\lambda) = F(x|\lambda) = \sum_{i=1}^n \lambda_i f_i(x) < \sum_{i=1}^n \lambda_i f_i(\hat{x}) = F(\hat{x}|\lambda),$$

which is a contradiction.

- **Case 3:** If $f_i(\hat{x}) < f_i(x)$ and $f_j(\hat{x}) > f_j(x)$, then

$$\begin{aligned} F^*(\hat{\lambda}) &= F(\hat{x}|\hat{\lambda}) + \varepsilon(f_i(\hat{x}) - f_j(\hat{x})) \geq F(x|\hat{\lambda}) \\ &= F(x|\lambda) + \varepsilon(f_i(x) - f_j(x)). \end{aligned}$$

But because $\varepsilon(f_i(x) - f_j(x)) > \varepsilon(f_i(\hat{x}) - f_j(\hat{x}))$, this implies

$$\begin{aligned} F(\hat{x}|\lambda) &\geq F(x|\hat{\lambda}) - \varepsilon(f_i(\hat{x}) - f_j(\hat{x})) \\ &= \underbrace{F(x|\lambda)}_{=F^*(\lambda)} + \underbrace{\varepsilon(f_i(x) - f_j(x)) - \varepsilon(f_i(\hat{x}) - f_j(\hat{x}))}_{>0} \\ &> F^*(\lambda), \end{aligned}$$

which is a contradiction.

Hence, $f_i(\hat{x}) \geq f_i(x)$ and $f_j(\hat{x}) \leq f_j(x)$, as claimed. $\square$

Proof of Lemma 1. Fix $\delta > 0$ and $i \in \mathcal{N}$, and let $(x, \hat{x}) \in \mathcal{X}(w) \times \mathcal{X}(\hat{w})$. Then

$$\begin{aligned} F^*(\hat{w}) - F^*(w) &= F(\hat{x}|\hat{w}) - F(x|w) \\ &= \underbrace{F(\hat{x}|\hat{w}) - F(x|\hat{w})}_{\geq 0} + \underbrace{F(x|\hat{w}) - F(x|w)}_{=\delta f_i(x)} \geq \delta f_i(x), \end{aligned}$$

because, by assumption, $\hat{w} = w + \delta e_i$. As the preceding inequality holds for any $x \in \mathcal{X}(w)$, it follows that

$$F^*(\hat{w}) - F^*(w) \geq \delta(\max_{x \in \mathcal{X}(w)} f_i(x)) > 0,$$

and thus $F^*(\hat{w}) > F^*(w)$, which completes the proof. $\square$

Proof of Proposition 2. We first show that, for any given decision, the performance ratio is quasiconcave in the weight. To that end, fix any feasible action $x \in \mathcal{X}$, and define

$$\mathcal{U}_c(x) = \{\lambda \in \Delta : \varphi(x|\lambda) \geq c\} \quad (\text{A.1})$$

as the set of (normalized) weights λ that yield a performance ratio $\varphi(x|\lambda)$ in Equation (3) at a level of at least $c \in [0, 1]$. In fact, the upper contour sets in Equation (A.1) are nested in their level, that is,

$$c \leq \hat{c} \Rightarrow \mathcal{U}_{\hat{c}}(x) \subseteq \mathcal{U}_c(x), \quad (\text{A.2})$$

for any $\hat{c} \in [0, 1]$. Furthermore, for sufficiently small values of c , we have $\mathcal{U}_c(x) = \Delta$.⁴⁶ Thus, we may choose $c \in [0, 1]$ such that $\mathcal{U}_c(x)$ is nonempty, and consider any two weights $\lambda, \hat{\lambda} \in \mathcal{U}_c(x)$. For any $\theta \in (0, 1)$, their convex combination $\lambda_\theta = \theta\lambda + (1 - \theta)\hat{\lambda}$ belongs to $\mathcal{U}_c(x)$ if and only if

$$\varphi(x|\lambda_\theta) = \frac{F(x|\lambda_\theta)}{F^*(\lambda_\theta)} \geq c. \quad (\text{A.3})$$

Because $\varphi(x|\lambda) \geq c$ and $\varphi(x|\hat{\lambda}) \geq c$, the definition of the weighted objective in Equation (1) implies

$$\begin{aligned} F(x|\lambda_\theta) &= F(x|\theta\lambda + (1 - \theta)\hat{\lambda}) = \theta F(x|\lambda) + (1 - \theta)F(x|\hat{\lambda}) \\ &\geq c(\theta F^*(\lambda) + (1 - \theta)F^*(\hat{\lambda})), \end{aligned} \quad (\text{A.4})$$

where we use that $\min\{F^*(\lambda), F^*(\hat{\lambda})\} > 0$ by the nontriviality condition (N). Moreover, the convex combination of the individually optimized objectives $F^*(\lambda)$ and $F^*(\hat{\lambda})$ (each obtained from possibly different decisions) cannot be smaller than the value of the jointly optimized objective (from a single decision):

$$\begin{aligned} \theta F^*(\lambda) + (1 - \theta)F^*(\hat{\lambda}) &= \max_{x, \hat{x} \in \mathcal{X}} \{\theta F(x|\lambda) + (1 - \theta)F(\hat{x}|\hat{\lambda})\} \\ &\geq \max_{x \in \mathcal{X}} \{\theta F(x|\lambda) + (1 - \theta)F(x|\hat{\lambda})\} \\ &= \max_{x \in \mathcal{X}} \{F(x|\theta\lambda + (1 - \theta)\hat{\lambda})\} = F^*(\lambda_\theta). \end{aligned} \quad (\text{A.5})$$

Therefore, as long as $c \geq 0$, the fact that $F^*(\lambda_\theta) > 0$, together with Equations (A.4) and (A.5), implies that Equation (A.3) holds. As an immediate consequence, it is $\lambda_\theta \in \mathcal{U}_c(x)$, which in turn means that $\mathcal{U}_c(x)$ is convex for all $c \in [0, 1]$ and all $x \in \mathcal{X}$, as claimed.

The convexity of the upper contour sets $\mathcal{U}_c(\mathbf{x})$ implies that $\varphi(\mathbf{x}|\boldsymbol{\lambda})$ is quasiconcave in $\boldsymbol{\lambda}$ (see, e.g., Arrow and Enthoven 1961, p. 780). Hence, the minimum of the performance ratio over all weights in Δ is attained on the boundary:

$$\min_{\boldsymbol{\lambda} \in \Delta} \varphi(\mathbf{x}|\boldsymbol{\lambda}) = \min_{\boldsymbol{\lambda} \in \partial \Delta} \varphi(\mathbf{x}|\boldsymbol{\lambda}).$$

The same quasiconcavity argument applies to any face of Δ , so the minimum must also lie on the boundary of each face. Repeating this argument recursively over edges and vertices, we conclude that the minimum is attained in the set $\Lambda = \{e_1, \dots, e_n\}$ of the n vertices of Δ . Thus, we obtain a form of “perfect complementarity,”

$$\begin{aligned} \rho(\mathbf{x}) &= \min_{\boldsymbol{\lambda} \in \{e_1, \dots, e_n\}} \left\{ \frac{F(\mathbf{x}|\boldsymbol{\lambda})}{F^*(\boldsymbol{\lambda})} \right\} = \min_{i \in \mathcal{N}} \left\{ \frac{F(\mathbf{x}|e_i)}{F^*(e_i)} \right\} = \min_{i \in \mathcal{N}} \left\{ \frac{f_i(\mathbf{x})}{f_i^*} \right\} \\ &= \min_{i \in \mathcal{N}} \phi_i(\mathbf{x}), \end{aligned}$$

as claimed in Equation (6), which concludes the proof. $\square$

Proof of Lemma 2. For any $\mathbf{x} \in \mathcal{X}$, Proposition 2 provides for a representation of the performance index, which can be rewritten in the form $\rho(\mathbf{x}) = \phi_{\iota(\mathbf{x})}$, where $\iota(\mathbf{x}) \in \mathcal{I}(\mathbf{x}) = \arg \min_{i \in \mathcal{N}} \phi_i(\mathbf{x})$. The set of pseudo-robust decisions Ψ in Equation (7) is nonempty by the extreme-value theorem (Rudin 1976, theorem 4.16, p. 89), and it is compact by the maximum theorem (Berge 1963, p. 116). These two theorems also guarantee that the correspondence $\mathcal{I}(\cdot)$ is upper semi-continuous and nonempty-valued. Consider now $\mathcal{P}$, which by Equations (9) and (10) can be written as

$$\mathcal{P} = \{\mathbf{x} \in \mathcal{X} : \nexists \mathbf{x}' \in \mathcal{X} \text{ s.t. } \mathbf{x}' \succ \mathbf{x}\} = \{\mathbf{x} \in \mathcal{X} : \mathbf{x}' \preceq \mathbf{x}, \mathbf{x}' \in \mathcal{X}\},$$

where $\preceq$ is the negation of $\succ$. By the continuity of f , the preference $\succ$ is continuous, so that $\mathcal{P}$ is compact. We now show that it must also be nonempty. For this, start with any $\xi^0 \in \Psi$: If $\xi^0 \in \mathcal{P}$, then immediately $\xi^0 \in \mathcal{R}$, so the claim holds. Otherwise, if $\xi^0 \notin \mathcal{P}$, then there exists $\xi^1 \in \mathcal{X}$ such that $\xi^1 \succ \xi^0$. By Equation (9) there exists $i \in \mathcal{N}$ such that $f_i(\xi^1) > f_i(\xi^0)$ and $f(\xi^1) \geq f(\xi^0)$. Equivalently, $\phi_i(\xi^1) > \phi_i(\xi^0)$ and $\phi_j(\xi^1) \geq \phi_j(\xi^0)$, for all $j \in \mathcal{N}$, so that $\rho(\xi^1) \geq \rho(\xi^0)$. Because by assumption $\xi^0 \in \Psi$, by Equation (7) it must be $\rho(\xi^1) \geq \rho(\xi^0) = \rho^* = \max \rho(\mathcal{X})$, so $\xi^1 \in \Psi$ as well. In this manner one can now construct a sequence of decisions ξ^k , for $k \in \{0, 1, \dots\}$, continuing until $\xi^k \in \mathcal{P}$ for some $k \geq 0$, which then establishes that $\Psi \cap \mathcal{P}$ is indeed nonempty. Failing that, by the Bolzano-Weierstrass theorem (Berge 1963, p. 67) the sequence $(\xi^k)_{k \geq 0}$ must have a convergent subsequence in the compact set $\mathcal{X}$. Because the convergence is monotone in the (continuous) preference, there can only be a single accumulation point, which is equal to the limit of the sequence: $\hat{\mathbf{x}} = \lim_{k \rightarrow \infty} \xi^k$. Given that the set of all pseudo-robust decisions Ψ is compact (i.e., in particular closed) as noted earlier, it is $\hat{\mathbf{x}} \in \Psi$. In addition, $\mathbf{x}' \preceq \hat{\mathbf{x}}$ for all $\mathbf{x}' \in \mathcal{X}$, as the upper and lower contour sets of a continuous preference are closed. But this means that $\hat{\mathbf{x}} \in \Psi \cap \mathcal{P}$, so the robust decision set $\mathcal{R}$ is both nonempty and compact, for it is the intersection of two compact sets. $\square$

Proof of Lemma 3. By Lemma 2, the robust decision set is nonempty and compact. Hence, by the extreme-value theorem there exists $\hat{\mathbf{x}}^* \in \mathcal{R} = \Psi \cap \mathcal{P}$, and Equation (7) yields $\rho(\mathcal{R}) = \rho(\Psi \cap \mathcal{P}) = \{\rho^*\} = \rho(\Psi)$, as well as $\max \rho(\mathcal{P}) = \rho^* = \max \rho(\mathcal{X})$, establishing the claim. $\square$

Proof of Lemma 4. (i) The claim follows by combining Equations (7) and (12). (ii) Fix $\varepsilon \in (0, 1]$, and—following the outcome-based logic discussed in Remark 9—consider a selection

$$\mathbf{y}_\varepsilon^* \in \mathcal{Y}_\varepsilon = \arg \max_{\mathbf{y} \in \mathcal{Y}} \left\{ \sum_{i \in \mathcal{N}} y_i : t_\varepsilon \leq y_1, \dots, t_\varepsilon \leq y_n, t_\varepsilon \in \mathcal{T}_\varepsilon \right\}.$$

We now show that $\mathbf{y}_\varepsilon^*$ is efficient (with respect to the coordinates of points in $\mathcal{Y}$).⁴⁷ Suppose it is not; then there exists a feasible $\hat{\mathbf{y}} > \mathbf{y}_\varepsilon^*$ which achieves a strictly greater payoff, a contradiction, which in turn establishes the claimed efficiency of $\mathbf{y}_\varepsilon^*$. Moreover, it is clear that $\mathcal{Y}_\varepsilon = \phi(\mathcal{R}_\varepsilon)$, so that

$$\mathcal{R}_\varepsilon = \phi^{-1}(\mathcal{Y}_\varepsilon) = \{\mathbf{x} \in \mathcal{X} : \phi(\mathbf{x}) \in \mathcal{Y}_\varepsilon\}.$$

As a result, $\mathcal{R}_\varepsilon \subset \mathcal{P}$, for all $\varepsilon \in (0, 1]$, as claimed. (iii) Let $\varepsilon \in (0, 1]$, and consider any $(\mathbf{x}, \hat{\mathbf{x}}) \in \mathcal{R}_\varepsilon \times \Psi$. If $\mathbf{x} \notin \Psi$, then

$$\rho = \min_{i \in \mathcal{N}} \phi_i(\mathbf{x}) < \min_{i \in \mathcal{N}} \phi_i(\hat{\mathbf{x}}) = \rho^* = \max \rho(\mathcal{X}).$$

On the other hand, $\mathbf{x} \in \mathcal{R}_\varepsilon$ implies that $\Phi_\varepsilon(\mathbf{x}) \geq \Phi_\varepsilon(\hat{\mathbf{x}})$. Let $\kappa = \sum_{i \in \mathcal{N}} (\phi_i(\mathbf{x}) - \phi_i(\hat{\mathbf{x}}))$ be the total coordinate-wise difference in performance. Then for any $\hat{\varepsilon} > 0$,

$$\Phi_\varepsilon(\mathbf{x}) \geq \Phi_\varepsilon(\hat{\mathbf{x}}) \Leftrightarrow \hat{\varepsilon} \cdot \kappa \geq \rho^* - \rho > 0 \Leftrightarrow \hat{\varepsilon} \geq \frac{\rho^* - \rho}{\kappa} > 0.$$

This establishes the claim in part (iii), for any $\varepsilon' \in (0, \min\{\varepsilon, (\rho^* - \rho)/\kappa\})$ and $\hat{\varepsilon} \in (0, \varepsilon')$. For $\hat{\varepsilon} = 0$, the claim follows from part (i). $\square$

Proof of Proposition 3. We first note that $\mathcal{R}_\varepsilon$ is upper semicontinuous in $\varepsilon \in [0, 1]$, and set $\mathcal{Q} = \varliminf_{\varepsilon \rightarrow 0^+} \mathcal{R}_\varepsilon$. The proof proceeds in two steps. We first show that $\mathcal{Q} \subset \mathcal{R}$ and then $\mathcal{Q} \neq \emptyset$.

Step 1: $\mathcal{Q} \subset \mathcal{R}$. By upper semicontinuity of $\mathcal{R}_\varepsilon$ we have that $\mathcal{Q} \subset \mathcal{R}_0 = \Psi$ (see, e.g., Aubin and Frankowska 1990, p. 41), taking into account that Ψ is compact; cf. Endnote 19 for Section 3.1. Because $\mathcal{R}_\varepsilon \subset \mathcal{P}$, for all $\varepsilon \in (0, 1]$, we further obtain $\mathcal{Q} \subset \text{cl } \mathcal{P}$. Assume that there is a point $\mathbf{x} \in \mathcal{Q} \setminus \mathcal{P}$, which means that $\mathbf{x} = \lim_{\varepsilon \rightarrow 0^+} \mathbf{x}_\varepsilon$, for some selection $\mathbf{x}_\varepsilon \in \mathcal{R}_\varepsilon$, where $\varepsilon \in (0, 1]$. Because $\mathbf{x} \notin \mathcal{P}$, there exists $\hat{\mathbf{x}} \in \mathcal{P}$ such that $\hat{\mathbf{x}} \succ \mathbf{x}$. That is, $\phi_j(\hat{\mathbf{x}}) > \phi_j(\mathbf{x})$ for some $j \in \mathcal{N}$, and $\phi(\hat{\mathbf{x}}) \geq \phi(\mathbf{x})$. If we denote by $A = (1/n) \sum_{i \in \mathcal{N}} \phi_i$ the average performance over all criteria, then $\mathbf{x} \notin \mathcal{P}$ implies that $A(\mathbf{x}) < A(\hat{\mathbf{x}})$. By the continuity of $A(\cdot)$ on $\mathcal{X}$, it is therefore $\lim_{\varepsilon \rightarrow 0^+} A(\mathbf{x}_\varepsilon) = A(\mathbf{x}) < A(\hat{\mathbf{x}})$. Hence, there exists a $\bar{\varepsilon} \in (0, 1)$ such that $A(\mathbf{x}_\varepsilon) < A(\hat{\mathbf{x}})$ for all $\varepsilon \in (0, \bar{\varepsilon})$. Thus, for any given $\varepsilon \in (0, \bar{\varepsilon})$:

$$\begin{aligned} \Phi_\varepsilon(\hat{\mathbf{x}}) - \Phi_\varepsilon(\mathbf{x}_\varepsilon) &= (1 - \varepsilon)(\rho(\hat{\mathbf{x}}) - \rho(\mathbf{x}_\varepsilon)) + \varepsilon(A(\hat{\mathbf{x}}) - A(\mathbf{x}_\varepsilon)) \\ &\geq \varepsilon(A(\hat{\mathbf{x}}) - A(\mathbf{x}_\varepsilon)) > 0, \end{aligned}$$

where we have taken into account that $\hat{\mathbf{x}} \in \Psi$, so $\rho(\hat{\mathbf{x}}) = \max \rho(\mathcal{X}) \geq \rho(\mathbf{x}_\varepsilon)$. But this means $\Phi_\varepsilon(\mathbf{x}_\varepsilon) < \Phi_\varepsilon(\hat{\mathbf{x}})$, contradicting $\mathbf{x}_\varepsilon \in \mathcal{R}_\varepsilon$. Hence, we have shown that necessarily $\mathcal{Q} \setminus \mathcal{P} = \emptyset$. Because we already know that $\mathcal{Q} \subset \Psi \cap (\text{cl } \mathcal{P})$, one obtains $\mathcal{Q} \subset \Psi \cap \mathcal{P} = \mathcal{R}$, as claimed.

Step 2: $\mathcal{Q} \neq \emptyset$. Consider a monotone sequence $(\varepsilon_k)_{k=0}^\infty \subset (0, 1)$ with $\varepsilon_{k+1} < \varepsilon_k$, for all $k \geq 0$, and such that $\lim_{k \rightarrow \infty} \varepsilon_k = 0$. Using the same selection $\mathbf{x}_\varepsilon$ as in Step 1, the sequence $(\mathbf{x}_{\varepsilon_k})_{k=0}^\infty \subset \mathcal{X}$ is a sequence of points contained in the compact set $\mathcal{X}$, so that by the Bolzano-Weierstrass theorem (Berge 1963, p. 67) there exists a convergent subsequence $(\mathbf{x}_{\varepsilon_{k_j}})_{j=0}^\infty \subset$

$(x_{\varepsilon_k})_{k=0}^\infty$ (with limit in $\mathcal{X}$). That is, there exists $q \in \mathcal{X}$, such that $q = \lim_{j \rightarrow \infty} x_{\varepsilon_{k_j}}$, and by the definition of $\mathcal{Q}$ we therefore obtain that $q \in \mathcal{Q}$ and thus, $\mathcal{Q} \neq \emptyset$, as stated in the result.

This concludes the proof. $\square$

Proof of Lemma 5. Continuity of the optimal ε -augmented performance index Φ_ε^* over $\varepsilon \in [0, 1]$ follows from the maximum theorem (Berge 1963, p. 116), given the continuity of Φ_ε in both arguments and the compactness of $\mathcal{X}$. To establish monotonicity, fix $\varepsilon', \varepsilon'' \in [0, 1]$ with $\varepsilon'' > \varepsilon'$. By optimality of $x_{\varepsilon''}$ for $\Phi_{\varepsilon''}$ and feasibility of $x_{\varepsilon'}$ at ε'' , we obtain

$$\Phi_{\varepsilon''}^* = (1 - \varepsilon'')\rho_{\varepsilon''} + \varepsilon''\mu_{\varepsilon''} \geq (1 - \varepsilon'')\rho_{\varepsilon'} + \varepsilon''\mu_{\varepsilon'} = \Phi_{\varepsilon''}(x_{\varepsilon'}).$$

Subtracting $\Phi_{\varepsilon'}^* = (1 - \varepsilon')\rho_{\varepsilon'} + \varepsilon'\mu_{\varepsilon'}$, we find

$$\Phi_{\varepsilon''}^* - \Phi_{\varepsilon'}^* \geq \Phi_{\varepsilon''}(x_{\varepsilon'}) - \Phi_{\varepsilon'}^* = (\varepsilon'' - \varepsilon')(\mu_{\varepsilon'} - \rho_{\varepsilon'}) \geq 0,$$

because $\mu_{\varepsilon'} \geq \rho_{\varepsilon'}$ by definition. This confirms that Φ_ε^* is non-decreasing on $[0, 1]$. To establish convexity, we show that regular increments in ε increase the corresponding differences of the value function. If $\tilde{\varepsilon} = (\varepsilon' + \varepsilon'')/2$ denotes the midpoint between ε' and ε'' , it is sufficient to show that the difference of successive differences,

$$\Delta = (\Phi_{\varepsilon''}^* - \Phi_{\tilde{\varepsilon}}^*) - (\Phi_{\tilde{\varepsilon}}^* - \Phi_{\varepsilon'}^*),$$

must be nonnegative. Indeed, setting $\delta = (\varepsilon'' - \varepsilon')/2 > 0$, it is

$$\Phi_{\varepsilon''}^* - \Phi_{\tilde{\varepsilon}}^* \geq \delta(\mu_{\tilde{\varepsilon}} - \rho_{\tilde{\varepsilon}}) \geq \Phi_{\tilde{\varepsilon}}^* - \Phi_{\varepsilon'}^*,$$

where the first inequality follows from $\Phi_{\varepsilon''}^* \geq \Phi_{\varepsilon''}(x_{\tilde{\varepsilon}})$, and the second from $\Phi_{\varepsilon'}^* \geq \Phi_{\varepsilon'}(x_{\tilde{\varepsilon}})$. Together, these imply $\Delta \geq 0$, which proves that Φ_ε^* is convex on $[0, 1]$. $\square$

Proof of Lemma 6. (i)/(iv) The function $\Phi_\varepsilon(\cdot)$ conforms to Equation (1) as a weighted objective comprising the two criteria $\rho(\cdot)$ and $\mu(\cdot)$. By Proposition 1, increasing $\varepsilon \in [0, 1]$, which shifts weight from the first criterion to the second, can only decrease the optimal value of the first and increase that of the second. Thus, ρ_ε is decreasing and μ_ε is increasing in $\varepsilon \in [0, 1]$, as claimed. (ii) Because $\Phi_0 = \rho$, we have that $\rho_0 = \max_{x \in \mathcal{X}} \Phi_0(x) = \rho^*$. Hence, by part (i) we have that $\rho^* \geq \rho_\varepsilon$ for all $\varepsilon \in [0, 1]$. (iii) Because ρ_ε is bounded above by ρ^* and nonincreasing in ε , the sequence $(\rho_\varepsilon)_{\varepsilon \in [0, 1]}$ converges as $\varepsilon \rightarrow 1^-$ by the monotone convergence theorem (see, e.g., Rudin 1976, theorem 3.14, p. 55). Moreover, because ρ^* is the least upper bound, the limit of ρ_ε as $\varepsilon \rightarrow 1^-$ must equal ρ^* . (v) This result follows from the monotonicity of Φ_ε^* established in Lemma 5. Because $\Phi_\varepsilon^* = (1 - \varepsilon)\rho_\varepsilon + \varepsilon\mu_\varepsilon$ is nondecreasing in ε , and $\rho_\varepsilon \leq \rho^*$, by parts (i) and (ii), it must be that $\mu_\varepsilon \geq \rho^*$ for all $\varepsilon \in [0, 1]$; otherwise, the convex combination would decrease, contradicting monotonicity. $\square$

Proof of Proposition 4. By Equations (15) and (16), along with Lemma 5, we have

$$0 \leq \psi_\varepsilon = \rho^* - \rho_\varepsilon = \Phi_0^* - \frac{\Phi_\varepsilon^* - \varepsilon\mu_\varepsilon}{1 - \varepsilon} \leq \Phi_0^* - \Phi_\varepsilon^* + \varepsilon\mu_\varepsilon \leq \varepsilon\mu_\varepsilon, \quad \varepsilon \in [0, 1].$$

Thus, given a desired approximation-error bound $\delta > 0$, for any $\hat{\varepsilon} \in (0, 1]$, it is

$$\varepsilon \leq \min\{\hat{\varepsilon}, \delta/\mu_{\hat{\varepsilon}}\} \Rightarrow \psi_\varepsilon \leq \delta.$$

Because, by Lemma 6(iv), the function $\hat{\varepsilon} \mapsto \delta/\mu_{\hat{\varepsilon}}$ is decreasing,

a simpler implication follows by choosing $\hat{\varepsilon} = 1$:

$$\varepsilon \leq \delta/\mu_1 \Rightarrow \psi_\varepsilon \leq \delta.$$

This establishes the claimed approximation guarantee. $\square$

Proof of Lemma 7. Fix any $\varepsilon \in [0, 1]$. By Equation (16), the difference between the upper and lower bounds in Equation (17) is given by

$$d_\varepsilon = \Phi_\varepsilon^* - \rho_\varepsilon = \varepsilon(\mu_\varepsilon - \rho_\varepsilon) \geq 0.$$

The midpoint estimator $\hat{\rho}_\varepsilon$ defined in Equation (18) is the arithmetic average of these bounds:

$$\hat{\rho}_\varepsilon = \frac{\Phi_\varepsilon^* + \rho_\varepsilon}{2}.$$

By construction, this implies

$$|\hat{\rho}_\varepsilon - \rho^*| \leq \frac{1}{2}(\Phi_\varepsilon^* - \rho_\varepsilon) = \frac{d_\varepsilon}{2}.$$

Therefore, $\hat{\rho}_\varepsilon$ approximates ρ^* to within $d_\varepsilon/2$, as claimed. $\square$

Proof of Lemma 8. Let $\Omega \subset \mathbb{R}_+^n \setminus \{0\}$ be a nonempty compact set, which is not reduced to the origin.

i. (a) Because Ω is nonempty, its radial projection $\pi(\Omega)$ is nonempty as well. Moreover, π is continuous on the compact set Ω , so $\pi(\Omega)$ is compact (Apostol 1974, theorem 4.25, p. 82). We now show that $\pi(\Omega) = \mathcal{C}(\Omega) \cap \Delta$. For the inclusion $\pi(\Omega) \subset \mathcal{C}(\Omega) \cap \Delta$, take any $\lambda \in \pi(\Omega)$. Then there exists $w \in \Omega$ such that $\pi(w) = \lambda$, i.e., $\lambda = w/\|w\|_1$. Because $\Omega \subset \mathbb{R}_+^n \setminus \{0\}$, it follows that for any $\alpha > 0$, we have $\alpha w \in \mathcal{C}(\Omega)$, and in particular, $\lambda = \alpha w$ for $\alpha = 1/\|w\|_1$. Hence $\lambda \in \mathcal{C}(\Omega) \cap \Delta$. Conversely, suppose $\hat{\lambda} \in \mathcal{C}(\Omega) \cap \Delta$. Then there exists $\hat{\alpha} > 0$ such that $\hat{w} = \hat{\alpha}\hat{\lambda} \in \Omega$. Applying the radial projection yields $\pi(\hat{w}) = \hat{w}/\|\hat{w}\|_1 = \hat{\lambda}$, because $\|\hat{\lambda}\|_1 = 1$. Therefore, $\hat{\lambda} \in \pi(\Omega)$, proving $\mathcal{C}(\Omega) \cap \Delta \subset \pi(\Omega)$. Together, we obtain $\pi(\Omega) = \mathcal{C}(\Omega) \cap \Delta$. (b) Because $\partial\Omega \subset \Omega$, we trivially have $\pi(\partial\Omega) \subset \pi(\Omega)$. To prove the reverse inclusion, let $\hat{\lambda} \in \pi(\Omega)$. Then $\hat{\lambda} = \pi(\hat{w})$ for some $\hat{w} = \hat{\alpha}\hat{\lambda} \in \Omega$, with $\hat{\alpha} > 0$. Because Ω is compact and $0 \notin \Omega$, the ray $\alpha\hat{\lambda}$ intersects Ω in a bounded interval of positive length. That is, there exist real numbers $\underline{\alpha}, \bar{\alpha}$, with $0 < \underline{\alpha} \leq \hat{\alpha} \leq \bar{\alpha} < \infty$, such that $\underline{\alpha}\hat{\lambda}, \bar{\alpha}\hat{\lambda} \in \partial\Omega$. Hence, $\hat{\lambda} = \pi(\underline{\alpha}\hat{\lambda}) \in \pi(\partial\Omega)$, so $\pi(\Omega) \subset \pi(\partial\Omega)$, completing the proof that $\pi(\Omega) = \pi(\partial\Omega)$.

ii. Let $w \in \text{int}(\Omega)$. Because $\Omega \subset \mathcal{C}(\Omega)$, it follows that $w \in \text{int}(\mathcal{C}(\Omega))$. Because Δ is a smooth manifold and the intersection of an open set with it is open (in the relative topology), we obtain that $\pi(\text{int}(\Omega)) = \text{int}(\mathcal{C}(\Omega)) \cap \Delta$ is open in Δ ,⁴⁸ and thus $\pi(w) \in \text{int}(\pi(\Omega))$. This proves that the radial projection maps interior points of Ω to interior points of $\pi(\Omega)$. It follows that if $\Omega' \subset \Omega$ is open in $\mathbb{R}_+^n$, then $\pi(\Omega')$ is open in Δ .

iii. Let $\Omega' \subset \Omega$ be nonempty. Then $\pi(\Omega') \subset \Delta$ is nonempty. Its closure is compact because Ω' is contained in the compact set Ω . Using continuity of π and part (i) (b) we have that

$$\partial\pi(\Omega') = \partial(\text{cl}(\pi(\Omega'))) = \partial\pi(\text{cl}(\Omega')) = \partial\pi(\partial(\text{cl}(\Omega'))) = \partial\pi(\partial\Omega').$$

Hence, the boundary of $\pi(\Omega')$ equals the boundary of the projection of the boundary of Ω' .

iv. Suppose Ω is convex. If $\pi(\Omega)$ were not convex, then there would exist $\lambda', \lambda'' \in \pi(\Omega)$ and $\mu \in (0, 1)$ such that

$\lambda = \mu\lambda' + (1 - \mu)\lambda'' \notin \pi(\Omega)$. Let $w' \in \pi^{-1}(\{\lambda'\})$, and $w'' \in \pi^{-1}(\{\lambda''\})$. Then for any $\alpha > 0$,

$$\alpha\lambda = \frac{\alpha\mu}{\|w'\|_1}w' + \frac{\alpha(1-\mu)}{\|w''\|_1}w''.$$

In particular, define

$$\alpha = \left(\frac{\mu}{\|w'\|_1} + \frac{1-\mu}{\|w''\|_1} \right)^{-1} > 0, \text{ and } \hat{\mu} = \frac{\alpha\mu}{\|w'\|_1} \in (0, 1).$$

Then

$$\alpha\lambda = \hat{\mu}w' + (1 - \hat{\mu})w'' \in \Omega,$$

by convexity of Ω . Hence, $\pi(\alpha\lambda) = \lambda \in \pi(\Omega)$, contradicting our assumption. Therefore, $\pi(\Omega)$ must be convex.

v. We now show that the image $\pi(\Omega')$ of a straight line segment $\Omega' \subset \Omega$ is a straight line segment in Δ , with the possibility of a degenerate case when the straight line segment is projected to a single point in Δ . Indeed,

$$\begin{aligned} \pi(\hat{\mu}w' + (1 - \hat{\mu})w'') &= \frac{\hat{\mu}w' + (1 - \hat{\mu})w''}{\|\hat{\mu}w' + (1 - \hat{\mu})w''\|_1} = \frac{\hat{\mu}\|w'\|_1\lambda' + (1 - \hat{\mu})\|w''\|_1\lambda''}{\|\hat{\mu}w' + (1 - \hat{\mu})w''\|_1} \\ &= \frac{\hat{\mu}\|w'\|_1}{\hat{\mu}\|w'\|_1 + (1 - \hat{\mu})\|w''\|_1}\lambda' + \left(1 - \frac{\hat{\mu}\|w'\|_1}{\hat{\mu}\|w'\|_1 + (1 - \hat{\mu})\|w''\|_1}\right)\lambda'' \\ &= \mu(\hat{\mu})\lambda' + (1 - \mu(\hat{\mu}))\lambda'', \end{aligned}$$

where $\lambda' = w'/\|w'\|_1$, $\lambda'' = w''/\|w''\|_1$, and

$$\mu(\hat{\mu}) = \frac{\hat{\mu}\|w'\|_1}{\hat{\mu}\|w'\|_1 + (1 - \hat{\mu})\|w''\|_1}, \quad \hat{\mu} \in [0, 1].$$

The latter function is continuously differentiable, with

$$\frac{d\mu(\hat{\mu})}{d\hat{\mu}} = \frac{\|w'\|_1\|w''\|_1}{(\hat{\mu}\|w'\|_1 + (1 - \hat{\mu})\|w''\|_1)^2} > 0, \quad \hat{\mu} \in [0, 1],$$

which implies that $\mu(\cdot)$ is strictly increasing on $[0, 1]$. Thus, provided that $\lambda' \neq \lambda''$, the radial projection of a straight line is a bijection, which also preserves the orientation of any straight path between two different points in Ω , as long as they do not map to the same point in Δ (for $\lambda' = \lambda''$).

vi. Consider first the case where $\text{co}(\Omega)$ is a bounded polyhedron (i.e., a polytope; see, e.g., Nemhauser and Wolsey 1999, definition 2.2, p. 86). By part (i)(b) it is $\pi(\text{co}(\Omega)) = \pi(\partial(\text{co}(\Omega)))$. By part (iv), the set $\pi(\text{co}(\Omega))$ is convex. Part (v) then guarantees that the straight lines between any two vertices of $\text{co}(\Omega)$ remain straight lines in their radial projection onto Δ , which implies it is enough to project the vertices of $\text{co}(\Omega)$ and then take the convex hull. Because the vertices are included in $\partial\Omega$, we therefore obtain that $\pi(\text{co}(\Omega)) = \text{co}(\pi(\partial\Omega))$ as claimed. If the bounded convex set $\text{co}(\Omega)$ is not a polytope, then it can be approximated arbitrarily closely by a polytope (see, e.g., Bronstein 2008), so that the claim—omitting some of the convergence-specific details—follows in the limit.

vii. Let $\Omega', \Omega'' \subset \mathbb{R}_+^n$ such that $\Omega = \Omega' \cup \Omega''$. Then $\mathcal{C}(\Omega' \cup \Omega'') = \mathcal{C}(\Omega') \cup \mathcal{C}(\Omega'')$, which implies the claim by part (i)(a) (cf. Endnote 48): $\pi(\Omega) = (\mathcal{C}(\Omega') \cap \Delta) \cup (\mathcal{C}(\Omega'') \cap \Delta) = \pi(\Omega') \cup \pi(\Omega'')$.

viii. Let $\Omega', \Omega'' \subset \mathbb{R}_+^n$ such that $\Omega = \Omega' \cap \Omega''$. Then, by part (i)(a) we have $\pi(\Omega) = \mathcal{C}(\Omega' \cap \Omega'') \cap \Delta$. If $\Omega = \Omega' \cap \Omega''$ for some $\Omega', \Omega'' \subset \mathbb{R}_+^n$, then $\pi(\Omega) = \mathcal{C}(\Omega' \cap \Omega'') \cap \Delta$, as claimed. $\square$

Proof of Proposition 5. Recall that the ambiguity set $\Omega \subset \mathbb{R}_+^n$ is compact, nonempty, and not reduced to the origin. For any feasible decision $x \in \mathcal{X}$, the Ω -conditioned performance index corresponds to its worst-case relative performance. Using the radial projection of Ω and invoking Lemma 8(i), we obtain:⁴⁹

$$\rho(x|\Omega) = \min_{w \in \Omega} \frac{F(x|w)}{F^*(w)} = \min_{\lambda \in \pi(\Omega)} \frac{F(x|\lambda)}{F^*(\lambda)} = \min_{\lambda \in \pi(\partial\Omega)} \frac{F(x|\lambda)}{F^*(\lambda)}.$$

As shown in the proof of Proposition 2, the performance ratio $\varphi(x|\lambda) = F(x|\lambda)/F^*(\lambda)$ is quasiconcave in λ for fixed $x \in \mathcal{X}$. Because $\pi(\Omega)$ is compact, the minimum is attained on the boundary $\partial\pi(\Omega)$, which equals $\pi(\partial\Omega)$ by Lemma 8(i)(b). Moreover, the quasiconcavity of $\varphi(x|\cdot)$ implies that all upper contour sets are convex, so the minimum is also attained on the boundary of the convex hull $\text{co}(\pi(\partial\Omega))$. Thus, if there exists an extremal base Λ such that $\text{co}(\Lambda) = \text{co}(\pi(\Omega))$, then the minimum is attained on the compact set Λ , as claimed. $\square$

Proof of Proposition 6. Fix any $\theta \in \Theta$, and consider the state-contingent Ω' -conditioned performance index,

$$\phi(x, \theta|\Omega') = \min_{\lambda' \in \pi(\Omega')} \frac{G(x, \theta|\lambda')}{G^*(\theta|\lambda')}, \quad x \in \mathcal{X}, \quad (\text{A.6})$$

which measures the performance of an action relative to all feasible weighted-sum scalarizations of the multicriteria decision problem, in state θ . Because Λ' is an extremal base of $\pi(\Omega')$, we have $\text{co}(\Lambda') = \text{co}(\pi(\Omega'))$. The quasiconcavity of $G(x, \theta|\lambda')/G^*(\theta|\lambda')$ in λ' for fixed (x, θ) , which obtains as in the proof of Proposition 2, ensures that the minimum in Equation (A.6) is attained at an extreme point of $\text{co}(\pi(\Omega'))$. Thus, as in Proposition 5, we obtain an equivalent representation,

$$\phi(x, \theta|\Omega') = \min_{\lambda' \in \Lambda'} \frac{G(x, \theta|\lambda')}{G^*(\theta|\lambda')}, \quad x \in \mathcal{X}. \quad (\text{A.7})$$

By Equation (27), the performance index is the minimum of the state-contingent Ω' -conditioned performance index in Equation (A.7) over all $\theta \in \Theta$, that is,

$$\rho(x|\Omega', \Theta) = \min_{\theta \in \Theta} \phi(x, \theta) = \min_{\theta \in \Theta} \left\{ \min_{\lambda' \in \Lambda'} \frac{G(x, \theta|\lambda')}{G^*(\theta|\lambda')} \right\}, \quad x \in \mathcal{X}. \quad (\text{A.8})$$

Thus, reversing the order of minimization in Equation (A.8) yields Equation (27), as claimed. $\square$

Proof of Corollary 1. Let $\Lambda' = \{\lambda'^{(1)}, \dots, \lambda'^{(n)}\}$ be the (smallest) extremal base of $\pi(\Omega')$ (or, equivalently, of $\text{co}(\pi(\Omega'))$), which is finite by hypothesis, with $n \geq 1$. Then Equation (27) becomes

$$\begin{aligned} \rho(x|\Omega', \Theta) &= \min_{\lambda' \in \Lambda'} \left\{ \min_{\theta \in \Theta} \frac{G(x, \theta|\lambda')}{G^*(\theta|\lambda')} \right\} = \min_{i \in \mathcal{N}} \left\{ \min_{\theta \in \Theta} \frac{G(x, \theta|\lambda'^{(i)})}{G^*(\theta|\lambda'^{(i)})} \right\} \\ &= \min_{i \in \mathcal{N}} \{f_i(x)\}, \quad x \in \mathcal{X}, \end{aligned} \quad (\text{A.9})$$

where we have set $f_i(x) = \min_{\theta \in \Theta} \{G(x, \theta | \lambda^{(i)}) / G^*(\theta | \lambda^{(i)})\}$, for all $(x, i) \in \mathcal{X} \times \mathcal{N}$. Because by definition $\max_{x \in \mathcal{X}} G(x, \theta | \lambda^{(i)}) = G^*(\theta | \lambda^{(i)})$, it is $f_i = \max_{x \in \mathcal{X}} f_i(x) = 1$, for all $i \in \mathcal{N}$, so that Equation (A.9) is in fact equivalent to Equation (28), which completes the proof. $\square$

Proof of Lemma 9. Fix $i, j, m \in \mathcal{N}$ (with $m \neq n$), $\lambda \in \Delta$, and $\delta > 0$. Then for any $y \in \mathcal{Y}$: If $\hat{y} = y + \delta e_i \in \mathcal{Y}$, then

$$\begin{aligned} H(y | \lambda) &= h^{-1} \left(\sum_{l=1}^n \lambda_l h(y_l) \right) \\ &= h^{-1} \left(\sum_{l=1}^n \lambda_l h(\hat{y}_l) - \lambda_i (h(\hat{y}_i) - h(y_i)) \right) \\ &\begin{cases} < H(\hat{y} | \lambda), & \text{if } \lambda_i > 0, \\ = H(\hat{y} | \lambda), & \text{if } \lambda_i = 0. \end{cases} \end{aligned}$$

This implies Axiom 1.⁵⁰ Consider now the case where $(\hat{y}, \hat{\lambda}) = (y, \lambda) + (y_i - y_j, \lambda_i - \lambda_j)(e_j - e_i)$. Then for $i = j$, Axiom 2 holds trivially. For $i \neq j$, it is

$$H(\hat{y} | \hat{\lambda}) = h^{-1} \left(\sum_{l \in \mathcal{N} \setminus \{i, j\}} \lambda_l h(y_l) + \lambda_i h(y_i) + \lambda_j h(y_j) \right) = H(y | \lambda),$$

which also yields Axiom 2. Because λ is by assumption normalized, we also obtain that Axiom 3 holds, as

$$H(\alpha \mathbf{1}_n) = h^{-1} \left(\sum_{l=1}^n \lambda_l h(\alpha) \right) = h^{-1}(h(\alpha)) = \alpha,$$

for any $\alpha > 0$. As far as associativity is concerned, let us assume that $(\lambda_1, \dots, \lambda_m) \neq \mathbf{0}$, and let

$$\alpha = H \left(\sum_{l=1}^m y_l e_l \middle| \frac{\sum_{l=1}^m \lambda_l e_l}{\sum_{l=1}^m \lambda_l} \right) = h^{-1} \left(\frac{\sum_{l=1}^m \lambda_l h(y_l)}{\sum_{l=1}^m \lambda_l} \right).$$

With this, one obtains

$$\begin{aligned} &H(\alpha \mathbf{1}_m, y_{m+1}, \dots, y_n | \lambda) \\ &= h^{-1} \left(h \circ h^{-1} \left(\frac{\sum_{l=1}^m \lambda_l h(y_l)}{\sum_{l=1}^m \lambda_l} \right) \cdot \left(\sum_{l=1}^m \lambda_l \right) + \sum_{l=m+1}^n \lambda_l h(y_l) \right) \\ &= H(y | \lambda), \end{aligned}$$

whence Axiom 4 must hold (taking into account that $h \circ h^{-1}(\cdot) = (\cdot)$). Finally, we observe that

$$H(y | e_i) = h^{-1} \circ h(y_i) = y_i,$$

so the coordinate-filter requirement, which constitutes Axiom 5, is also satisfied. This concludes the proof. $\square$

Endnotes

¹ Practical examples for multicriteria decision making are legion; see, for example, the collections of case studies by Berbel et al. (2018) in agriculture, Masri et al. (2018) in financial decision making, and Ravindran (2016) for supply chain management, to name just a few.

² Extant ESG metrics differ widely. In their approach to “quantifying the impact of impact investing,” Lo and Zhang (2024) remain agnostic about the impact factors to be used, taking them as given and thus keeping at bay the difficulties of attributing weights to different ESG criteria and of dealing with this model uncertainty.

³ Lancaster (1966) already noted that a product can be viewed as a bundle of its attributes, with consumer valuations often empirically assessed using conjoint analysis (see, e.g., Green and Srinivasan 1990).

⁴ Example 9 in Section 3.6 discusses robust allocation in a two-agent exchange economy using the proposed framework.

⁵ This is notwithstanding the “norm equivalence” in a finite-dimensional Euclidean space, in the sense that any norm $\|\cdot\|$ can be bracketed by any other norm $\|\cdot\|$, so $\chi_1 \|\cdot\| \leq \|\cdot\| \leq \chi_2 \|\cdot\|$ for suitable scalars $\chi_1, \chi_2 > 0$.

⁶ For the computation of Pareto sets, see Kung et al. (1975), Gabow et al. (1984), and Bentley et al. (1993).

⁷ See Yamamoto (2002) for details about maximizing a function on a Pareto-efficient set in a polyhedral setting.

⁸ A lexicographic evaluation of criteria based on perceived importance may justify an “elimination by aspects” (Tversky 1972), or more nuanced partial elimination heuristics using “attribute filters” (Kimya 2018).

⁹ Numerous ad hoc methods for determining weights exist, for example, in reliability engineering based on standards (Jiang and Chen 2020).

¹⁰ In his “Discussion on Making Things Equal,” Zhuang Zhou points out that “[t]here is nothing in the world bigger than the tip of an autumn hair, and Mount T’ai is tiny” (Watson 1968, p. 44), where Mount Tai is the highest point in the Shandong province of China.

¹¹ For example, getting \$100 in a month, in addition to repayment of the invested principal, would be attractive if obtained by investing a principal of \$1 today, but not if it required investing a principal of \$10 million today.

¹² A point $x \in \mathcal{X}$ is said to be *isolated* if there exists an open set $\mathcal{O} \subset \mathbb{R}^m$ such that $\mathcal{O} \cap \mathcal{X} = \{x\}$.

¹³ By the extreme-value theorem (Rudin 1976, theorem 4.16, p. 89), the minimum of a continuous function on a compact set exists and is finite.

¹⁴ It is $\mathbb{R}_+^n \setminus \{0\} = \bigcup_{\alpha > 0} \alpha \Delta$.

¹⁵ More precisely, as shown in the proof of Lemma 1, we have $F^*(\hat{w}) \geq F^*(w) + \delta(\max f_i(\mathcal{X}(w)))$.

¹⁶ This “complete ignorance” (or *full ambiguity*) is relaxed in Section 3.7 where we allow the decision maker to face *limited* ambiguity.

¹⁷ Because $F^*(\lambda) \geq \min F^*(\Delta) > 0$, the function $\varphi: \mathcal{X} \times \Delta \rightarrow [0, 1]$ is well defined.

¹⁸ Leontief (1941) employed this aggregation in fixed proportions as a simplification for his analysis of a larger economy.

¹⁹ By the extreme-value theorem, Ψ is nonempty, and by the maximum theorem, it is compact.

²⁰ Given two vectors $z, \hat{z} \in \mathbb{R}^n$, where $z = (z_1, \dots, z_n)$ and $\hat{z} = (\hat{z}_1, \dots, \hat{z}_n)$, the standard vector inequalities are defined as follows: (i) $z \leq \hat{z} \Leftrightarrow z_i \leq \hat{z}_i, \forall i \in \mathcal{N}$; (ii) $z \ll \hat{z} \Leftrightarrow z_i < \hat{z}_i, \forall i \in \mathcal{N}$; (iii) $z < \hat{z} \Leftrightarrow (z \leq \hat{z} \text{ and } \exists j \in \mathcal{N} \text{ such that } z_j < \hat{z}_j)$. Here, $\mathcal{N} = \{1, \dots, n\}$, with $n \geq 2$.

²¹ The relevant lower (set) limit is given by $\lim_{\varepsilon \rightarrow 0^+} \mathcal{R}_\varepsilon = \{x \in \mathcal{X} : \lim_{\varepsilon \rightarrow 0^+} (\inf_{x' \in \mathcal{R}_\varepsilon} \|x - x'\|) = 0\}$ (see, e.g., Aubin and Frankowska 1990, definition 1.4.6, p. 41). Given a sequence $(\varepsilon_k)_{k=1}^\infty \subset (0, 1]$ with $\lim_{k \rightarrow \infty} \varepsilon_k = 0$, this lower limit contains the accumulation points of any sequence $(x_k)_{k=1}^\infty$ with elements $x_k \in \mathcal{R}_{\varepsilon_k}$ for all $k \geq 1$.

²² Note that: $\Phi_\varepsilon(2, \frac{3}{2}) - \max_{x_1 \in [0, 1]} \left\{ \frac{(1-\varepsilon)x_1}{2} + \frac{\varepsilon}{2} \left(\frac{x_1}{2} + \frac{2+\sqrt{1-x_1}}{3} \right) \right\} = \frac{\varepsilon}{3} \frac{1-(7/12)\varepsilon}{2-\varepsilon} > 0$, for all $\varepsilon \in (0, 1]$.

²³ The larger (path-connected, compact) action set $\mathcal{X}' = \mathcal{X} \cup \{(2, \frac{3}{2}) : \zeta \in [0, 1]\}$ leaves results unchanged.

²⁴ It is $\frac{1}{24} (16 - 3\varepsilon^2) / (2 - \varepsilon) - \Phi_\varepsilon(3, \frac{1}{2}) = \frac{1}{24} (8 - 16\varepsilon + 7\varepsilon^2) / (2 - \varepsilon) > 0$ if and only if $\varepsilon < \frac{4}{7} (2 - \frac{1}{\sqrt{2}})$.

²⁵ Normalizing the maximum criteria to the same score (e.g., $c \in \{1, 10, 100\}$) is natural in many applications.

²⁶ The term “robustness equivalence principle” is analogous to the well-known “certainty equivalence principle,” for example, in linear-quadratic optimal control problems (Bertsekas 1995, p. 23), where it is optimal (i.e., maximizing the expected value of a quadratic objective functional) to replace uncertain parameters by their means.

²⁷ For $\alpha = \beta$, Equations (22) and (23) imply the unique robust allocation $\hat{x} = (1/2, 1/2)$, achieving $\rho^* = 1/2$.

²⁸ The converse of part (iv) does not hold. That is, if Ω is nonconvex, then $\pi(\Omega)$ may still be convex.

²⁹ The analysis in Example 10 yields an extreme counterexample for $\Omega = \{w\} = \Omega' \cap \Omega''$, with $\Omega' = [\underline{w}, w]$, $\Omega'' = [w, \bar{w}]$, and $\underline{w} \ll w \ll \bar{w}$. Indeed, for $\|\underline{w}\|_1 \rightarrow 0^+$ one obtains $\text{int}(\Delta) \subset \pi(\Omega')$, so $\pi(\Omega') \cap \pi(\Omega'') = \pi(\Omega'')$, whereas $\pi(\Omega) = \pi(\{w\})$ is merely a singleton.

³⁰ The underlying justification is Minkowski’s theorem, which states that any compact convex set is equal to the convex hull of its extreme points; see, for example, Rockafellar (1970, corollary 18.5.1, p. 167) or Schneider (2014, corollary 1.4.5, p. 17).

³¹ In general, the smallest Δ may not be finite (e.g., if $\Omega \subset \mathbb{R}_+^3$ is a unit ball, then the smallest $\Delta = \partial\pi(\Omega)$ has a continuum of elements).

³² For analogous comments about the assumed continuity of the criteria in the action, see Remark 1 (i) in Section 2.

³³ For example, in the case where $\Theta = \{(1, 0), (0, 1)\}$ for $l = n' = 2$, observing $\theta_1 = 1$ implies $\theta_2 = 0$ (and vice versa).

³⁴ For any $m \in \mathcal{N}$, we define $\mathbf{1}_m = (1, \dots, 1) \in \mathbb{R}_+^m$.

³⁵ It is $(\sum_{i=1}^m \lambda_i e_i) / (\sum_{i=1}^m \lambda_i) = \pi(\sum_{i=1}^m \lambda_i e_i)$, with the radial projection $\pi(\cdot)$, from $\mathbb{R}_+^n \setminus \{0\}$ onto Δ ; see Section 3.8.

³⁶ Kolmogorov (1930) showed that if Axioms 1–4 are satisfied for $\lambda = \mathbf{1}_n/n$ (cf. Endnote 34), then there exists a (continuous) strictly monotonic function $h: \mathcal{D} \rightarrow \mathbb{R}$, defined on a suitable domain $\mathcal{D} \subset \mathbb{R}$, such that

$$H(y|\mathbf{1}_n/n) = h^{-1}(\sum_{i=1}^n (1/n)h(y_i)), \quad y \in \mathcal{Y}.$$

Nagumo (1930) obtained a similar result, and the preceding “quasiarithmetic mean” is therefore also known as the “Kolmogorov-Nagumo mean.”

³⁷ Homogeneity is not required by Axioms 1–5, because relatively robust decisions (cf. Proposition 2) are invariant under rescaling.

³⁸ The gradient of LSE is the softmax function (Boltzmann 1871; Gibbs 1902, chapter XIV), widely used in machine learning (Bridle 1990).

³⁹ The purpose of restricting all criteria to strictly positive values was merely to skip checking the nontriviality condition (N) in Section 2.1 which is satisfied here (e.g., $x_d = 2$). Hence, setting $s_{13} = 0$ poses no problem.

⁴⁰ Under some mild additional assumptions (which we omit), the law of large numbers would guarantee convergence of the average scores to the true criterion-function values, at least when the number of available score observations for each realized action (i.e., each element of $\mathcal{X}$) goes to infinity, assuming that decisions can be properly recognized, as may naturally be the case on a sufficiently discrete underlying action set $\mathcal{X}$.

⁴¹ Nathan and Lord (1983) propose “knowledge,” “delivery,” “relevance,” “interpersonal,” and “organization” as criteria for the quality of a university lecturer; see also DeNisi and Murphy (2017) for a survey.

⁴² This is not the only “energy trilemma”; Tilman et al. (2009) discuss a food-energy-environment trilemma in the context of biofuels.

⁴³ Klarman et al. (1968) recognized the need for quality-of-life considerations in cost-effective treatment options (for chronic renal disease).

⁴⁴ Cutler and Richardson (1997), Buxton (2005), and Newhouse (2021) discuss QALY for assessing cost-effectiveness in healthcare.

⁴⁵ Baker and Freeland (1975) provide an overview of early scoring models. Stewart (1991) proposes a multicriteria decision support system for project selection using essentially equal weights for different objectives (akin to the Laplacian approach; cf. Section 4.1). Finally, Hall et al. (2015) suggest a project evaluation with an “underperformance riskiness index” relying on (typically unavailable) past statistical data.

⁴⁶ Indeed, because $\varphi(x|\lambda) \geq \min_{i \in \mathcal{N}} \{f_i(x)/F(x_i)\} = \underline{c}(x) \in [0, 1]$, for all $\lambda \in \Delta$, Equation (A.2) implies that $c \leq \underline{c}(x) \Rightarrow \mathcal{U}_c(x) = \Delta$.

⁴⁷ Efficiency in the outcome set corresponds directly to efficiency in the action set (cf. Section 3.2), because the latter is evaluated via a criterion vector that produces points in the outcome set.

⁴⁸ The representation $\pi(\Omega) = \mathcal{C}(\Omega) \cap \Delta$ was given in part (i) for a compact set Ω , but it also applies to any subset of Ω .

⁴⁹ By part (a) of Lemma 8 (i), $\pi(\Omega)$ is nonempty and compact, so the minimum with respect to λ must be attained by the extreme-value theorem.

⁵⁰ When h is increasing, the argument of h^{-1} cannot become smaller by dropping the term $\lambda_i(h(\hat{y}_i) - h(y_i))$. Conversely, when h is decreasing, the argument of h^{-1} weakly decreases, so the value of that inverse cannot go down.

References

- Apostol TM (1974) *Mathematical Analysis*, 2nd ed. (Addison-Wesley, Reading, MA).
- Arrow KJ, Enthoven AC (1961) Quasi-concave programming. *Econometrica* 29(4):779–800.
- Aubin J-P, Frankowska H (1990) *Set-Valued Analysis* (Birkhäuser, Boston).
- Baker N, Freeland J (1975) Recent advances in R&D benefit measurement and project selection methods. *Management Sci.* 21(10):1164–1175.
- Ben-David S, Borodin A (1994) A new measure for the study of on-line algorithms. *Algorithmica* 11(1):73–91.
- Ben-Tal A, den Hertog D, Waegenaere AD, Melenberg B, Rennen G (2013) Robust solutions of optimization problems affected by uncertain probabilities. *Management Sci.* 59(2):341–357.
- Bentley JL, Clarkson KL, Levine DB (1993) Fast linear expected-time algorithms for computing maxima and convex hulls. *Algorithmica* 9(2):168–183.
- Berbel J, Bourmaris T, Manos B, Matsatsinis N, Viaggi D, eds. (2018) *Multicriteria Analysis in Agriculture: Current Trends and Recent Applications* (Springer, New York).
- Berge C (1963) *Topological Spaces* (Oliver and Boyd, Edinburgh, UK).
- Bertsekas DP (1995) *Dynamic Programming and Optimal Control*, vol. 1 (Athena Scientific, Belmont, MA).
- Blanchet J, Murthy K (2019) Quantifying distributional model risk via optimal transport. *Math. Oper. Res.* 44(2):565–600.
- Boltzmann L (1871) Über das Wärmegleichgewicht zwischen mehratomigen Gasmolekülen. *Sitzungsberichte der Mathematisch-Naturwissenschaftlichen Classe der Kaiserlichen Akademie der Wissenschaften, Wien* 63(3):397–418.
- Bridle JS (1990) Probabilistic interpretation of feedforward classification network outputs, with relationships to statistical pattern recognition. Fogelman Soulié F, Hérault J, eds. *Neurocomputing: Algorithms, Architectures and Applications*, NATO ASI Series, vol. 68 (Springer, New York), 227–236.
- Bronstein EM (2008) Approximation of convex sets by polytopes. *J. Math. Sci.* 153(6):727–762. [Translation of Russian original, which appeared in *Sovremennaya Matematika / Fundamental'nyye Napravleniya* (Contemporary Mathematics / Fundamental Directions), Vol. 22, Geometry, 2007.]

- Buxton MJ (2005) How much are health-care systems prepared to pay to produce QALY? *Eur. J. Health Econom.* 6(4):285–287.
- Charnes A, Cooper WW (1961) *Management Models and Industrial Applications of Linear Programming*, vol. 1 (Wiley, New York).
- Charnes A, Cooper WW, Ferguson RO (1955) Optimal estimation of executive compensation by linear programming. *Management Sci.* 1(2):138–151.
- Cournot A (1838) *Recherches sur les Principes Mathématiques de la Théorie des Richesses* (Hachette, Paris).
- Cutler DM, Richardson E (1997) Measuring the health of the U.S. population. *Brookings Papers Econom. Activity Microeconom.* 28(1997): 217–282.
- Delage E, Ye Y (2010) Distributionally robust optimization under moment uncertainty with application to data-driven problems. *Oper. Res.* 58(3):595–612.
- DeNisi AS, Murphy KR (2017) Performance appraisal and performance management: 100 years of progress? *J. Appl. Psych.* 102(3):421–433.
- Edgeworth FY (1881) *Mathematical Psychics* (C. Kegan Paul, London).
- Edgeworth FY (1897) La teoria pura del monopolio. *Giornale degli Economisti* 15(Anno 8):13–31.
- Ehrgott M (2010) *Multicriteria Optimization*, 2nd ed. (Springer, Berlin).
- Fechner GT (1860) *Elemente der Psychophysik (Part II)* (Breitkopf & Härtel, Leipzig, Germany).
- Gabow HN, Bentley JL, Tarjan RE (1984) Scaling and related techniques for geometry problems. *Proc. 16th Annual ACM Sympos. Theory Comput.* (Association for Computing Machinery, New York), 135–143.
- Gibbs JW (1902) *Elementary Principles in Statistical Mechanics* (General Publishing, Toronto).
- Goel A, Meyerson A, Weber TA (2009) Fair welfare maximization. *Econom. Theory* 41(3):465–494.
- Green PE, Srinivasan V (1990) Conjoint analysis in marketing: New developments with implications for research and practice. *J. Marketing* 54(4):3–19.
- Hall NG, Long DZ, Qi J, Sim M (2015) Managing underperformance risk in project portfolio selection. *Oper. Res.* 63(3):660–675.
- Han J, Weber TA (2023) Price discrimination with robust beliefs. *Eur. J. Oper. Res.* 306(2):795–809.
- Hardy GH, Littlewood JE, Pólya G (1934) *Inequalities* (Cambridge University Press, Cambridge, UK).
- Harsanyi JC (1955) Cardinal welfare, individualistic ethics, and interpersonal comparisons of utility. *J. Political Econom.* 63(4):309–321.
- Hawkins TR, Gausen OM, Strømman AH (2012) Environmental impacts of hybrid and electric vehicles: A review. *Internat. J. Life Cycle Assessment* 17(8):997–1014.
- Jiang R, Chen Z (2020) A standard-based approach for multi-criteria performance evaluation of engineered systems. *Reliability Engrg. System Safety* 202(107001):1–12.
- Kahneman D, Tversky A (1979) Prospect theory: An analysis of decision under risk. *Econometrica* 47(2):263–291.
- Kahneman D, Tversky A (1984) Choices, values, and frames. *Amer. Psychologist* 39(4):341–350.
- Keeney RL, Raiffa H (1993) *Decisions with Multiple Objectives* (Cambridge University Press, Cambridge, UK).
- Kimya M (2018) Choice, consideration sets, and attribute filters. *Amer. Econom. J. Microeconom.* 10(4):223–247.
- Klarman HE, Francis JOS, Rosenthal GD (1968) Cost effectiveness analysis applied to the treatment of chronic renal disease. *Medical Care* 6(1):48–54.
- Kolmogorov AN (1930) Sur la notion de la moyenne. *Atti della Accademia Nazionale dei Lincei, Classe di Scienze Fisiche, Matematiche e Naturali Rendiconti* 12(9):388–391.
- Kouvelis P, Yu G (1997) *Robust Discrete Optimization and Its Applications* (Springer, New York).
- Kung H, Luccio F, Preparata F (1975) On finding the maxima of a set of vectors. *J. ACM* 22(4):469–476.
- Lancaster KJ (1966) A new approach to consumer theory. *J. Political Econom.* 74(2):132–157.
- Leontief WW (1941) *The Structure of American Economy 1919–1929: An Empirical Application of Equilibrium Analysis* (Harvard University Press, Cambridge, MA).
- Lo AW, Zhang R (2024) Quantifying the impact of impact investing. *Management Sci.* 70(10):7161–7186.
- Loewenstein GF, Thompson L, Bazerman MH (1989) Social utility and decision making in interpersonal contexts. *J. Personality Soc. Psych.* 57(3):426–441.
- Marshall A (1890) *Principles of Economics* (Macmillan, London).
- Mas-Colell A, Whinston MD, Green JR (1995) *Microeconomic Theory* (Oxford University Press, Oxford, UK).
- Masri H, Pérez-Gladish B, Zopounidis C, eds. (2018) *Financial Decision Aid Using Multiple Criteria: Recent Models and Applications* (Springer, New York).
- Miyamoto JM, Wakker PP, Bleichrodt H, Peters HJM (1998) The zero-condition: A simplifying assumption in QALY measurement and multiattribute utility. *Management Sci.* 44(6):839–849.
- Mohajerin Esfahani P, Kuhn D (2018) Data-driven distributionally robust optimization using the Wasserstein metric: Performance guarantees and tractable reformulations. *Math. Programming* 171(1–2):115–166.
- Nagumo M (1930) Über eine Klasse der Mittelwerte. *Japanese J. Math.* 7:71–79.
- Nathan BR, Lord RG (1983) Cognitive categorization and dimensional schemata: A process approach to the study of halo in performance ratings. *J. Appl. Psychol.* 68(1):102–114.
- Nemhauser GL, Wolsey LA (1999) *Integer and Combinatorial Optimization* (Wiley Interscience, New York).
- Newhouse JP (2021) An ounce of prevention. *J. Econom. Perspect.* 35(2):101–118.
- Pareto V (1894) Il massimo di utilità dato dalla libera concorrenza. *Giornale degli Economisti* 9(2):48–66.
- Pareto V (1897) *Cours d'Économie Politique*, vol. 2 (F. Rouge, Lausanne, Switzerland).
- Pareto V (1906) *Manuale di Economia Politica* (Società Editrice Libreria, Milan).
- Pliskin JS (1974) The management of patients with end-stage renal failure: A decision-theoretic approach. Doctoral thesis, Harvard University, Boston.
- Pliskin JS, Beck CH (1976) A health index for patient selection: A value function approach with application to chronic renal failure patients. *Management Sci.* 22(9):1009–1021.
- Ravindran AR, ed. (2016) *Multiple Criteria Decision Making in Supply Chain Management* (CRC Press, New York).
- Rockafellar RT (1970) *Convex Analysis* (Princeton University Press, Princeton, NJ).
- Rudin W (1976) *Principles of Mathematical Analysis*, 3rd ed. (McGraw-Hill, New York).
- Savage LJ (1951) The theory of statistical decision. *J. Amer. Statist. Assoc.* 46(253):55–67.
- Schneider R (2014) *Convex Bodies: The Brunn-Minkowski Theory*, 2nd ed. (Cambridge University Press, Cambridge, UK).
- Sleator DD, Tarjan RE (1985) Amortized efficiency of list update and paging rules. *Comm. ACM* 28(2):202–208.
- Steuer RE (1986) *Multiple Criteria Optimization: Theory, Computation, and Application* (Wiley, New York).
- Stewart TJ (1991) A multi-criteria decision support system for R&D project selection. *J. Oper. Res. Soc.* 42(1):17–26.
- Tilman D, Socolow R, Foley JA, Hill J, Larson E, Lynd L, Pacala S, et al. (2009) Beneficial biofuels: The food, energy, and environment trilemma. *Science* 325(5938):270–271.
- Tversky A (1972) Elimination by aspects: A theory of choice. *Psych. Rev.* 79(4):281–299.
- Villani C (2008) *Optimal Transport: Old and New*, Grundlehren der Mathematischen Wissenschaften, vol. 338 (Springer, Berlin).

- Wald A (1945) Statistical decision functions which minimize the maximum risk. *Ann. Math.* 46(2):265–280.
- Wang T, Fu Y (2020) Constructing composite indicators with individual judgements and best-worst method. *Soc. Indicators Res.* 149(1):1–14.
- Watson B (1968) *The Complete Works of Chuang Tzu* (Columbia University Press, New York).
- Weber EH (1846) Der Tastsinn und das Gemeingefühl. Wagner R, ed. *Handwörterbuch der Physiologie mit Rücksicht auf die physiologische Pathologie*, vol. 3, section 2 (Vieweg, Braunschweig, Germany), 481–588.
- Weber TA (2014) On the (non-)equivalence of IRR and NPV. *J. Math. Econom.* 52:25–39.
- Weber TA (2023) Relatively robust decisions. *Theory Decision* 94(1): 35–62.
- Weber TA (2024) Optimal depth of discharge for electric batteries with robust capacity-shrinkage estimator. *Proc. 2024 4th Internat. Conf. Smart Grid Renewable Energy (SGRE)* (IEEE, New York), 1–5.
- Weber TA (2025) Monopoly pricing with unknown demand. *Scandinavian J. Econom.* 127(1):235–285.
- World Energy Council (2024) World Energy Trilemma 2024: Evolving with resilience and justice. Technical report, World Energy Council, London.
- Yager RR (1988) On ordered weighted averaging aggregation operators in multi-criteria decision making. *IEEE Trans. Systems Man Cybernetics* 18(1):183–190.
- Yamamoto Y (2002) Optimization over the efficient set: Overview. *J. Global Optim.* 22(1–4):285–317.
- Zarghami M, Szidarovszky F (2010) On the relation between compromise programming and ordered weighted averaging operator. *Inform. Sci.* 180(11):2239–2248.
- Zeleny M (1974) A concept of compromise solutions and the method of the displaced ideal. *Comput. Oper. Res.* 1(3):479–496.