

18 *Beyond Money: Multi-Criteria Decision Analysis*

After this chapter you will be able to:

1. explain why monetizing every criterion and deciding by feel both hide the judgment in a decision
2. compute the weighted total $V_j = \sum_k w_k s_{kj}$ of the additive value model, after checking that the weights sum to exactly 1 and that a present worth of cost has been normalized to a 0–100 score
3. derive the flip weight at which two alternatives tie when one weight is swept and the other weights are held in proportion
4. build criterion weights from 1–9 pairwise judgments by row geometric means, and check them with the consistency ratio $CR = CI/RI$
5. diagnose the three abuses of weighted scoring: double counting, weighting after seeing the scores, and mixed scales
6. choose among options that deliver the same kind of unpriced outcome by comparing annual equivalent cost per unit of outcome

Every chapter so far has ended in a single unit: SAR at time zero, SAR per year, a percentage rate. chapter 17 stretched that unit as far as it defensibly goes, pricing carbon and end-of-life flows into the ledger. This chapter faces what remains. Real siting, procurement, and design decisions carry criteria that resist a price: water security, community acceptance, safety margin, expansion room. Two easy escapes present themselves, and this chapter rejects both. One escape is to monetize everything, inventing a SAR value for community trust with false precision; the other is to decide by feel, letting the loudest voice in the room weight the criteria silently. The middle path is *multi-criteria decision analysis* (MCDA):

keep the money criterion exactly as the earlier chapters computed it, put the unpriceable criteria beside it as scores, state the weights in public, and add. The additive model is developed first, then the sensitivity of its verdict to the weights, then the analytic hierarchy process (AHP) for building weights from pairwise judgments, then the standard abuses, and finally cost-effectiveness analysis, the bridge to the public-sector methods of chapter 8.

When money alone cannot rank the sites.

Q1. Two concentrator sites have PW of cost 210 million and 240 million SAR. The cheaper site sits on a stressed aquifer; the dearer one has secure water. Does present worth alone pick the site?

Q2. An analyst uses weights 0.5, 0.3, and 0.1 on three criteria. What is wrong before any scoring starts?

Q3. In a pairwise comparison, an engineer judges criterion A three times as important as B, and B twice as important as C. If the judgments are perfectly consistent, how important is A relative to C?

Answers.

Q1. No. PW ranks the sites on cost only. Water security is a real criterion with no defensible market price, so the decision needs a method that carries both criteria at once, with an explicit weight between them.

Q2. The weights sum to 0.9, not 1. Weighted totals then understate every alternative, and any comparison against an absolute threshold (“proceed if the score exceeds 70”) is meaningless. Renormalize so the weights sum to exactly 1.

Q3. $3 \times 2 = 6$ times as important. Real judgment matrices are never perfectly consistent; AHP measures how far they deviate, and a large deviation means the weights should not be trusted.

Multi-criteria decision analysis (MCDA). A family of methods for ranking alternatives on several criteria at once, with explicit scores and explicit weights, when no single unit can carry the whole decision.

18.1 When Money Alone Cannot Decide

Money is the best-behaved criterion an engineer will ever meet: it is cardinal, it aggregates across time through discounting, and two analysts with the same data will compute the same PW. When every consequence of a decision is a cash flow,

the earlier chapters are sufficient and this chapter is unnecessary. The trouble starts when consequences refuse the unit. What is the SAR value of a site the local community welcomes rather than tolerates? Of a plant layout that can double later? A century of practice supplies two tempting shortcuts.

The first shortcut is to *monetize everything*. Assign community acceptance a value of 8 million SAR, add it to the PW, and proceed as before. The appeal is that the old machinery keeps working; the flaw is that the 8 million is not an estimate of anything. It is a preference wearing the costume of a measurement, and once inside the PW it becomes indistinguishable from the audited capital numbers beside it. Where a defensible price exists, as for carbon in chapter 17, pricing is right; where none exists, inventing one launders judgment into arithmetic.

The second shortcut is to *decide by feel*: circulate the PW table, discuss around it, and let the committee's sense of the intangibles settle the choice. The appeal is honesty, since no fake numbers are created; the flaw is that the weighting still happens, silently, inside whoever speaks last and loudest, and it can never be audited, repeated, or defended to a regulator.

Insight. Both shortcuts fail the same test: they hide the judgment. The monetizer hides it inside an invented price; the committee hides it inside an unrecorded discussion. MCDA's entire contribution is to drag the judgment into daylight as two visible lists, the scores and the weights, so that anyone who disagrees can point at the exact number they dispute and rerun the addition. The method does not remove subjectivity; it *files* it where auditors can find it.

18.2 The Additive Value Model

The workhorse of practical MCDA is the additive value model, often called a SMART-style model (simple multi-attribute rating technique). For alternative j with score s_{kj} on criterion k and weights w_k :

$$V_j = \sum_k w_k s_{kj}, \quad \sum_k w_k = 1, \quad 0 \leq s_{kj} \leq 100. \quad (18.1)$$

The alternative with the largest V_j wins. The recipe has four steps: choose non-overlapping criteria; score every alternative on every criterion on the same 0–100 scale; choose weights summing to 1; multiply and add.

Criterion. A dimension on which alternatives are compared: cost, water security, community acceptance. Criteria must not overlap, or the overlapping concern is counted twice.

Score. A 0–100 rating of one alternative on one criterion, with 0 for the worst plausible outcome and 100 for the best. Scores come from data where data exist and from expert rating where they do not.

Weight. The stated importance of a criterion. Weights are nonnegative and **must sum to exactly 1**.

Money is not abandoned; it becomes the first criterion, and it enters through the front door of this book. Compute the PW of cost of each alternative exactly as in chapter 5, then convert to a score by linear normalization between the worst and best alternatives:

In practice. The non-cost scores are engineering outputs, not opinions with units. A water-security score comes from hydrogeology: aquifer test yields, drawdown projections, allocation seniority. A reliability score comes from failure-rate analysis. The weighted sum is only as good as the engineering studies feeding its columns, exactly as a PW is only as good as its cash-flow estimates.

$$s_{\text{cost},j} = 100 \times \frac{PW_{\text{worst}} - PW_j}{PW_{\text{worst}} - PW_{\text{best}}}, \tag{18.2}$$

so the cheapest alternative scores 100 and the dearest scores 0.

Example 18.1 — Siting a copper concentrator: three sites, four criteria

Najd Copper must site its new concentrator. Engineering studies produced a 20-year PW of cost for each site (chapter 5 machinery, MARR 10%) and 0–100 scores on three unpriced criteria:

- **Site A (wadi):** PW of cost 210 million SAR; water security 30 (stressed alluvial aquifer); community acceptance 55; expansion potential 45 (confined valley floor).
- **Site B (coastal):** PW of cost 240 million SAR; water security 95 (desalination intake); community acceptance 60; expansion potential 85.
- **Site C (plateau):** PW of cost 260 million SAR; water security 60 (deep wellfield); community acceptance 90 (far from settlements); expansion potential 70.

The steering committee adopts weights: cost 0.40, water security 0.25, community acceptance 0.20, expansion potential 0.15. Rank the sites.

Solution.

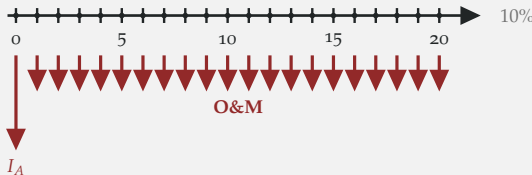


Figure 18.1. The cost criterion is ordinary engineering economy: each site’s 20-year cost stream is collapsed to a PW before it is scored. Site A’s diagram is shown; B and C have the same shape with their own numbers.

Step 1: score the cost criterion.

Normalize the PWs with Equation (18.2), worst 260, best 210:

$$s_A = 100, \quad s_B = 100 \times \frac{260 - 240}{260 - 210} = 40, \quad s_C = 0.$$

Step 2: assemble the decision matrix and check the weights.

The weights sum to $0.40 + 0.25 + 0.20 + 0.15 = 1.00$.

Criterion	weight w_k	Site A	Site B	Site C
Cost (normalized PW)	0.40	100	40	0
Water security	0.25	30	95	60
Community acceptance	0.20	55	60	90
Expansion potential	0.15	45	85	70
Weighted total V_j	1.00	65.25	64.50	43.50

Table 18.1. The concentrator-siting decision matrix. Every judgment in the analysis is visible in this table: the scores in the body, the weights in the second column, and the verdict in the last row.

Step 3: multiply and add.

$$V_A = 0.40(100) + 0.25(30) + 0.20(55) + 0.15(45)$$

$$= 40 + 7.5 + 11 + 6.75 = \boxed{65.25},$$

$$V_B = 0.40(40) + 0.25(95) + 0.20(60) + 0.15(85)$$

$$= 16 + 23.75 + 12 + 12.75 = \boxed{64.50},$$

$$V_C = 0.40(0) + 0.25(60) + 0.20(90) + 0.15(70)$$

$$= 0 + 15 + 18 + 10.5 = \boxed{43.50}.$$

Site A ranks first, by 0.75 points over Site B, with Site C well behind. Decision sentence: on the committee's stated weights, choose Site A, but the margin between A and B is less than one point on a 100-point scale, so the choice must be stress-tested against the weights before anyone signs, which is the subject of the next section.

Insight. What the additive model assumes. Adding weighted scores assumes the criteria trade off at constant rates: one point of water security is worth the same number of total points whether the site is parched or soaked. When a criterion has a cliff, for example a water level below which the plant simply cannot run, no weight is large enough to express it. Handle cliffs as *constraints*: screen out infeasible alternatives first, and score only the survivors. A weighted sum is for choosing among acceptable options, not for making unacceptable ones look acceptable.

Flip weight. The weight value at which two alternatives' weighted totals tie. It is the break-even analysis of chapter 12 with a weight, rather than a price or a volume, as the swept input.

18.3 Weight Sensitivity: The Flip Weight

Weights are the softest numbers in the model, so the verdict must be tested against them, exactly as chapter 12 tested a PW against its softest cash-flow estimates. The one-way recipe transfers directly: sweep one weight, hold the *relative* proportions of the other weights fixed (so the set still sums to 1), and find where the leaders tie.

Example 18.2 — The flip weight for the concentrator sites

In Example 18.1, Site A beat Site B by 0.75 points at a cost weight of 0.40. Sweep the cost weight w , keeping the other three weights in their original proportions 0.25 : 0.20 : 0.15, and find the flip weight between A and B.

Solution.

Step 1: collapse the non-cost criteria.

When cost has weight w , the other criteria share $1 - w$ in proportions 0.25/0.60, 0.20/0.60, 0.15/0.60. Each site's non-cost block is then a fixed average:

$$\bar{s}_A = \frac{0.25(30) + 0.20(55) + 0.15(45)}{0.60} = \frac{25.25}{0.60} = 42.08,$$

$$\bar{s}_B = \frac{0.25(95) + 0.20(60) + 0.15(85)}{0.60} = \frac{48.50}{0.60} = 80.83,$$

so $V_A(w) = 100w + 42.08(1 - w)$ and $V_B(w) = 40w + 80.83(1 - w)$.

Step 2: set them equal.

$$100w + 42.08(1 - w) = 40w + 80.83(1 - w) \implies 98.75w = 38.75,$$

$$w^* = \boxed{0.392}.$$

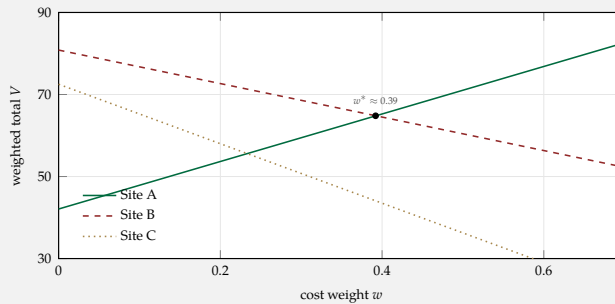


Figure 18.2. Weighted totals as the cost weight sweeps, other weights held in proportion. A and B cross at the flip weight; C never leads.

Step 3: read the decision regions.

Above a cost weight of 0.392, Site A leads; below it, Site B leads; Site C trails everywhere. The committee's chosen 0.40 sits *two hundredths* above the flip. Decision sentence: the ranking is not a fact about the sites, it is a fact about the weights; the committee should either defend cost mattering at least 39% of the total, or accept that A and B are, on present information, tied, and send the water and community studies back for sharper scores.

Insight. The flip weight converts an argument nobody can win into a question somebody can answer. Committee members will never agree that cost's weight is exactly 0.40 rather than 0.35; they do not need to. They need only agree on which side of 0.392 it falls, and the entire disagreement below that threshold is harmless, because it cannot change the decision. This is the same economy of effort that break-even analysis bought in chapter 12: precision is spent only where the verdict turns.

Reference. Saaty (1980) introduced AHP. The case study at the end of this chapter applies it to a solar-siting problem in Saudi Arabia.

Analytic hierarchy process (AHP). A method that derives criterion weights from a matrix of pairwise importance judgments on a 1–9 scale, and measures how self-contradictory the judgments are.

18.4 Where Do Weights Come From? A Brief, Honest AHP

Handing a committee a blank line labeled “weight” invites anchoring and horse-trading. The analytic hierarchy process asks an easier question many times instead: *of these two criteria, which matters more, and by how much?* Judgments use Saaty’s scale: 1 means equal importance, 3 moderate, 5 strong, 7 very strong, 9 extreme, with even values in between. The judgments fill a square matrix A where entry $a_{k\ell}$ is the judged importance of criterion k over criterion ℓ , and reciprocity is enforced: $a_{\ell k} = 1/a_{k\ell}$.

A practical way to extract weights, and the one used here, is the *row geometric mean*: multiply each row’s entries, take the n -th root, and normalize the results to sum to 1.

Example 18.3 — AHP weights for a pump procurement

A maintenance team must weight three criteria for choosing a slurry pump: life-cycle cost (C), reliability (R), and maintainability (M). Their pairwise judgments: C is moderately favored over R ($a_{CR} = 2$), C is favored over M ($a_{CM} = 4$), and R is moderately favored over M ($a_{RM} = 3$). Derive the weights and check consistency.

Solution.

Step 1: the judgment matrix.

Fill in the reciprocals:

$$A = \begin{pmatrix} 1 & 2 & 4 \\ 1/2 & 1 & 3 \\ 1/4 & 1/3 & 1 \end{pmatrix}.$$

Step 2: row geometric means.

$$r_C = (1 \times 2 \times 4)^{1/3} = 8^{1/3} = 2.000,$$

$$r_R = (0.5 \times 1 \times 3)^{1/3} = 1.5^{1/3} = 1.145,$$

$$r_M = (0.25 \times 0.333 \times 1)^{1/3} = 0.0833^{1/3} = 0.437.$$

Step 3: normalize.

The sum is $2.000 + 1.145 + 0.437 = 3.582$:

$$w_C = \boxed{0.558}, \quad w_R = \boxed{0.320}, \quad w_M = \boxed{0.122}.$$

Step 4: the consistency check.

Perfect consistency would require $a_{CM} = a_{CR} \times a_{RM} = 2 \times 3 = 6$, but the team said 4, so the judgments contradict themselves mildly. AHP quantifies this: compute Aw , divide each component by the corresponding w , and average to get $\lambda_{\max} = 3.018$; the consistency index is $CI = (\lambda_{\max} - n)/(n - 1) = 0.018/2 = 0.009$, and dividing by the random index for $n = 3$ ($RI = 0.58$) gives the consistency ratio

$$CR = \frac{0.009}{0.58} = \boxed{0.016},$$

far below the conventional 0.10 threshold. The judgments are usable, and the team can proceed with weights (0.558, 0.320, 0.122).

Insight. An honest reading of AHP. What it buys: pairwise questions are cognitively easier than direct weight-setting, the reciprocal structure forces each judgment to be made deliberately, and the consistency ratio catches self-contradiction that a table of directly chosen weights would hide forever. What it does not buy: objectivity. The weights are still one team's judgments, dressed in decimals; a different team produces different weights, legitimately. The correct use is to derive the weights, report the matrix and the CR alongside them, and then run the flip-weight analysis of section 18.3 anyway. The wrong use, common in consulting decks, is to present three-decimal AHP weights as if the decimals were data.

18.5 Common Abuses of Weighted Scoring

The additive model is easy to run and therefore easy to run dishonestly. Three abuses account for most of the damage in practice.

Insight. Double counting. If “capital cost” and “affordability” both appear as criteria, the cost concern votes twice, and whichever alternative is cheap wins twice with one virtue. The same happens with “reliability” and “downtime,” or “water security” and “drought risk.” The test is pairwise: for every two criteria, ask what information one adds that the other lacks; if the answer is nothing, merge them and re-normalize the weights. Criteria must partition the decision’s concerns, not restate them.

Insight. Weighting after seeing the scores. Run the scores first, notice that the favored alternative trails, then nudge the weights until it leads: this converts the model into a rationalization engine while leaving its paperwork immaculate, which is what makes it dangerous. The defense is procedural, not mathematical: fix the weights, in writing, before the scoring matrix is assembled, and treat any later weight change as a formal revision that requires a reason unrelated to the current ranking. The flip-weight analysis helps here too: a decision memo that already discloses “A leads for any cost weight above 0.39” leaves nowhere for a quiet nudge to hide.

In practice. The abuse that audits catch least often is score drift: a score quietly re-rated after a stakeholder meeting, with no new data behind it. Version the scoring matrix like a drawing: every score change carries a date, an author, and a cited study. An unversioned scoring sheet is as untrustworthy as an unversioned P&ID.

Cost-effectiveness analysis (CEA). Comparing alternatives that deliver the same *kind* of non-monetary outcome by the ratio of cost to outcome: SAR per tonne of dust avoided, SAR per injury prevented.

A third abuse is quieter: mixing scales. If three criteria are scored 0–100 and a fourth is entered on a 0–10 scale, the fourth criterion’s influence is silently divided by ten, whatever its stated weight. All scores must live on the same scale before the weights mean anything.

18.6 Cost-Effectiveness Analysis

Between full monetization and full scoring sits a large, practical middle case: the benefit is real, unpriced, and *one-dimensional*. Nobody defends a SAR value for a tonne of dust kept off a village, yet three dust-control options can still be ranked without one, because each delivers the same kind of outcome. Compute each option’s annual equivalent cost with the tools of chapter 6, divide by the annual outcome, and compare the ratios.

Example 18.4 — Three ways to keep dust down

A mine must suppress dust on its haul roads; the MARR is 10%. Three

options deliver the following (engineering-measured) reductions:

- **W (water trucks):** trucks cost 800,000 SAR, last 5 years, salvage 100,000 SAR; operations 350,000 SAR per year; avoids 60 tonnes of dust per year.
- **P (polymer sealant):** application plant costs 1,500,000 SAR, lasts 5 years, no salvage; operations 260,000 SAR per year; avoids 85 tonnes per year.
- **V (paving key segments):** costs 4,200,000 SAR, lasts 15 years, no salvage; upkeep 60,000 SAR per year; avoids 110 tonnes per year.

Rank the options by cost-effectiveness.

Solution.

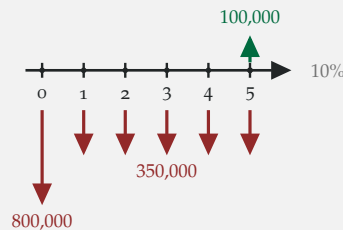


Figure 18.3. Option W's cash flows; each option is collapsed to an annual equivalent cost before the ratios are taken.

Step 1: annual equivalent costs.

With $(A/P, 10\%, 5) = 0.26380$ and $(A/P, 10\%, 15) = 0.13147$:

$$\begin{aligned} AEC_W &= (800,000 - 100,000)(0.26380) + 0.10(100,000) + 350,000 \\ &= 184,660 + 10,000 + 350,000 = 544,660 \text{ SAR/yr,} \end{aligned}$$

$$AEC_P = 1,500,000(0.26380) + 260,000 = 655,700 \text{ SAR/yr,}$$

$$AEC_V = 4,200,000(0.13147) + 60,000 = 612,174 \text{ SAR/yr.}$$

Step 2: divide by the outcome.

$$\begin{aligned} \frac{544,660}{60} &= \boxed{9,078 \text{ SAR/tonne}}, & \frac{655,700}{85} &= \boxed{7,714 \text{ SAR/tonne}}, \\ \frac{612,174}{110} &= \boxed{5,565 \text{ SAR/tonne}}. \end{aligned}$$

Step 3: read the ranking.

Paving buys dust reduction cheapest, at about 5,600 SAR per tonne, even though its AEC is neither the smallest nor its price tag modest; the water trucks, the cheapest option to start, are the most expensive per tonne delivered. Decision sentence: if the mine funds one option, fund paving; note that the ranking required no SAR value for a tonne of dust, only the engineering measurement of tonnes.

Insight. What CEA can and cannot conclude. It ranks options per unit of the same outcome, and it exposes dominated options instantly. What it cannot say is whether *any* option is worth doing: 5,565 SAR per tonne is the best available price for dust reduction, but only a value per tonne, however derived, can say whether the mine should buy at that price or how much reduction to buy. Setting such values for public decisions, and doing it with announced rules rather than invented numbers, is benefit–cost analysis, the subject of chapter 8.

18.7 Case Study from the Literature

MCDA is standard practice in energy and infrastructure siting, and the published studies use exactly the machinery of this chapter. Mohieldeen (2026), in an open-access *Scientific Reports* paper, maps the suitability of land around Makkah’s Holy Sites for utility-scale solar photovoltaic plants, the motivation being the 500–600 MW peak power demand of the Hajj. The study combines a geographic information system (GIS) with AHP: twelve criteria in four groups (climatic, topographic, environmental, and economic/infrastructure) are weighted by expert pairwise comparison, every 30-meter cell of the study area is scored on each criterion, and the weighted sum of Equation (18.1) is computed for the whole map at once. The AHP weights: solar radiation 21.5% (the largest), road-network proximity 17.5%, aspect 14.4%, slope 12.2%, land use 6.8%, wind speed 6.4%, precipitation 5.6%, settlement proximity 4.4%, humidity 4.3%, elevation 2.6%, power-line proximity 2.5%, and temperature 1.9%. The paper reports a consistency index of 0.036 against a random index of 1.54 for twelve criteria, and a consistency ratio of 0.024.

Reference. Mohieldeen (2026), *Scientific Reports*, 16. Open access; full reference at the end of the chapter.

Example 18.5 — Case study: checking the Makkah solar-siting model

Using only the published numbers of Mohieldeen (2026): (a) verify that the twelve criterion weights form a proper weight set; (b) verify the reported consistency ratio; (c) the paper classifies 110 km² as “Most Suitable” with a PV yield of 1,830 kWh per kWp-year, and applies a packing density of 30 MW per km²; check its claim that this land alone can cover the Hajj peak demand.

Solution.***(a) The weight set.***

Sum the twelve published weights:

$$21.5 + 17.5 + 14.4 + 12.2 + 6.8 + 6.4 + 5.6 + 4.4 + 4.3 + 2.6 + 2.5 + 1.9 = 100.1\%,$$

which is 1.00 within rounding of the published two-significant-figure weights: a proper weight set. (Its four group totals, 39.7% climatic, 29.2% topographic, 24.4% economic/infrastructure, and 6.8% environmental, likewise sum to 100.1%.)

(b) The consistency ratio.

With the published $CI = 0.036$ and $RI = 1.54$:

$$CR = \frac{CI}{RI} = \frac{0.036}{1.54} = 0.0234 \approx \boxed{0.024 \text{ (as published)}},$$

well below the 0.10 threshold, so the twelve-criterion judgment matrix is usable by the standard of section 18.4.

(c) The capacity claim.

Installed capacity on the best land:

$$110 \text{ km}^2 \times 30 \text{ MW/km}^2 = 3,300 \text{ MWp.}$$

Annual energy at the published yield:

$$3,300 \text{ MWp} \times 1,830 \text{ kWh/kWp} = 6.04 \text{ TWh per year,}$$

an average power of $6.04 \text{ TWh}/8,760 \text{ h} \approx \boxed{690 \text{ MW}}$, which exceeds the 500–600 MW Hajj peak. The paper’s arithmetic checks out; the qualitative

caveat, which the paper itself discusses, is that an average is not a peak, and evening pilgrimage load still needs storage or backup, the same day-block/night-block distinction as Example 17.2.

What did the researcher conclude? Mohieldeen (2026) finds that 10.38% of the study area is “Most Suitable” and a further 10.87% “High” suitability, concentrated in flat, sun-rich land northeast of the Holy Sites, near roads and a substation; 55.17% of the area is unsuitable or low-suitability, much of it on steep terrain or within buffers protecting the pilgrimage zones. The recommendation is a phased build-out starting with a 50–100 MW pilot plant in the northeastern zone. Two features of the study deserve a student’s attention. First, the money criteria (road-network proximity at 17.5%, power-line proximity at 2.5%) sit *inside* the weighted model rather than outside it, and the low power-line weight encodes a judgment that new lines can be built if a site is otherwise excellent. Second, the expert weights are disclosed in full, with their consistency ratio, so a reader who disagrees, for instance about temperature deserving only 1.9%, can rerun the map with different weights. That is MCDA practiced as this chapter teaches it: judgment made public, then arithmetic.

18.8 Chapter Summary

- Use MCDA when a decision carries criteria that resist defensible monetization. Reject both shortcuts: **monetizing everything** launders judgment into arithmetic; **deciding by feel** hides the weighting where it cannot be audited.
- The **additive value model**: $V_j = \sum_k w_k s_{kj}$ with weights summing to exactly 1 and all scores on one 0–100 scale. Cost enters as a criterion through the **normalized PW**, $s = 100 (PW_{\text{worst}} - PW) / (PW_{\text{worst}} - PW_{\text{best}})$.
- Criteria with cliffs are **constraints**: screen infeasible alternatives out first; the weighted sum chooses among survivors.
- The **flip weight** is break-even analysis (chapter 12) applied to a weight: sweep one weight with the others held in proportion, and find where the leaders tie. Committees need only agree on which side of the flip the true weight falls.

- **AHP** derives weights from 1–9 pairwise judgments: row geometric means, normalized. The **consistency ratio** $CR = CI/RI$ with $CI = (\lambda_{\max} - n)/(n - 1)$ flags self-contradiction; $CR \leq 0.10$ is the working threshold. AHP disciplines judgment; it does not create objectivity.
- The standard **abuses**: double counting (overlapping criteria), weighting after seeing the scores, and mixed scales. The defenses: pairwise overlap checks, weights fixed in writing before scoring, one scale for all scores.
- **Cost-effectiveness analysis** ranks options delivering the same kind of outcome by $AEC/\text{outcome}$ (SAR per tonne, SAR per injury avoided). It finds the cheapest way to buy an outcome but cannot say whether to buy at all; that question belongs to benefit–cost analysis (chapter 8).

18.9 Exercises

These exercises follow the chapter’s order: additive value models and score normalization, then flip weights, AHP weights with their consistency ratio, and cost-effectiveness ratios. The last four problems are longer: two exam-style sitting and procurement decisions, then two finished analyses to audit for double counting, inflated weight sets, and incoherent pairwise judgments.

Additive value models

Exercise 18.1. Two haul-road routes are scored on three criteria with weights 0.5 (cost), 0.3 (safety), 0.2 (dust exposure to the camp). Route N scores 80, 60, 40; Route S scores 60, 90, 70. Which route wins, and by how much?

Weight sensitivity

Exercise 18.2. Two flotation-reagent suppliers are scored on cost and on a combined non-cost block (quality, logistics, support). Supplier A: cost 100, non-cost 40. Supplier B: cost 30, non-cost 85. The cost weight is w and the non-cost block carries $1 - w$. The procurement team proposes $w = 0.5$. Find each supplier’s total at $w = 0.5$, find the flip weight, and write the decision sentence.

Deriving weights with AHP

Exercise 18.3. For a conveyor-belt supplier evaluation, a team judges: cost is moderately more important than reliability ($a_{CR} = 3$), cost is strongly more important than delivery time ($a_{CD} = 5$), and reliability is mildly more important than delivery time ($a_{RD} = 2$). Build the judgment matrix, derive the weights by row geometric means, and comment on consistency using a_{CD} against $a_{CR} \times a_{RD}$.

Exercise 18.4. An analyst enters criterion scores for four criteria, three on a 0–100 scale and one, by habit, on a 0–10 scale, then applies weights (0.4, 0.3, 0.2, 0.1) that sum correctly to 1. What happens to the fourth criterion’s influence on the ranking?

- (A) Nothing; the weight controls its influence.
- (B) Its influence is roughly ten times too small.
- (C) Its influence is roughly ten times too large.
- (D) The totals become invalid only if the weights also change.

Cost-effectiveness

Exercise 18.5. Three options reduce carbon monoxide exposure in an underground workshop. Option F (better fans): 240,000 SAR installed, 8-year life, no salvage, 30,000 SAR per year to run, cuts average exposure for 45 workers below the action level. Option S (source capture at 3 welding bays): 390,000 SAR installed, 8-year life, no salvage, 18,000 SAR per year, protects 60 workers. Option R (rotating schedules, no capital): 95,000 SAR per year in scheduling and relief labor, protects 25 workers. At $i = 10\%$, rank the options by cost per worker protected per year.

Exam-style problems

The following problems are written in the Major Exam format: a scenario, lettered parts with point values, and full work expected. Check that weights sum to 1 before computing, show the weighted-sum table, and end with a decision sentence.

Exercise 18.6. Siting a tailings facility at Wadi Phosphate. Three candidate sites have been screened as feasible. The 25-year PWs of cost are T1: 96 million, T2: 132 million, T3: 176 million SAR. Engineering studies give scores (0–100) on three further criteria:

	T1	T2	T3
Seepage risk (higher = safer)	50	85	95
Community acceptance	60	75	90
Closure and after-care	40	80	95

The board's weights are: cost 0.45, seepage 0.25, community 0.20, closure 0.10.

- (a) (10 pts) Normalize the cost PWs to 0–100 scores, verify the weight set, and compute the three weighted totals. (b) (6 pts) Sweep the cost weight with the other weights in proportion, and find the flip weight between the two leading sites. (c) (4 pts) Write the recommendation in one sentence, including the sensitivity finding.

Exercise 18.7. Choosing a dewatering pump by AHP. A team must choose between two pumps using three criteria: life-cycle cost (C), reliability (R), and maintainability (M). Their pairwise judgments: $a_{CR} = 2$, $a_{CM} = 5$, $a_{RM} = 3$. The engineering scores (0–100) are:

	C	R	M
Pump P1	60	90	70
Pump P2	90	55	65

- (a) (8 pts) Derive the criterion weights by row geometric means. (b) (6 pts) Compute $\lambda_{\max}$, the consistency index, and the consistency ratio ($RI = 0.58$ for $n = 3$), and state whether the judgments are usable. (c) (6 pts) Score the pumps and write the decision sentence.

Assessing outside recommendations

The two problems below present finished analyses produced by others. In each, the arithmetic shown is correct; the flaws are in the method. For each problem, (a) decide whether to accept the recommendation, (b) identify each error precisely, and (c) produce the corrected numbers.

Exercise 18.8. A consultant's scoring memo with too much weight. A consultant evaluates two water-treatment plants for a mine and submits the matrix below, recommending Plant 1 and noting that "its total of 81.0 also clears the client's 75-point approval threshold."

Weights: capital cost 0.35; affordability 0.20; reliability 0.25; water recovery 0.20; environmental compliance 0.15 (total 1.15, which slightly overweights and is conservative). Plant 1 scores: 95, 90, 50, 45, 55. Total: $0.35(95) + 0.20(90) + 0.25(50) + 0.20(45) + 0.15(55) = 81.0$. Plant 2 scores: 50, 55, 92, 88, 75. Total: 80.35. Recommendation: select Plant 1.

(a) (4 pts) Should the firm accept the recommendation? (b) (8 pts) Identify each error precisely. (c) (8 pts) Correct the model and produce the corrected ranking.

Exercise 18.9. An AHP table quoted to three decimals. A vendor's siting study reports criterion weights "derived by AHP" as $w_1 = 0.390$, $w_2 = 0.349$, $w_3 = 0.261$, quoted to three decimals, from the judgment matrix below.

$$A = \begin{pmatrix} 1 & 5 & 1/3 \\ 1/5 & 1 & 6 \\ 3 & 1/6 & 1 \end{pmatrix}$$

The geometric-mean weights are 0.390, 0.349, and 0.261. These weights were applied to the score matrix to obtain the final ranking.

(a) (4 pts) Verify that the stated weights do follow from the matrix. (b) (10 pts) Compute $\lambda_{\max}$, CI , and CR ($RI = 0.58$), and assess the study. (c) (6 pts) State what the vendor should be asked to do before the ranking is used.

References

- Mohieldeen, Y. (2026). GIS-based AHP multi-criteria mapping of potential solar PV power plant development: A case study in the vicinity of Holy Sites, Saudi Arabia. *Scientific Reports*, 16. <https://www.nature.com/articles/s41598-026-46353-9>
- Park, C. S. (2012). *Fundamentals of Engineering Economics* (3rd ed.). Pearson.
- Saaty, T. L. (1980). *The Analytic Hierarchy Process: Planning, Priority Setting, Resource Allocation*. McGraw-Hill.