

Marketing Analytics as Decision Support

Framing Decisions, Stakeholders, and the Analytic Specification

Dr. Jose Mendoza, Academic Director and Clinical Associate Professor

Version 1.0 · July 2026

Except where otherwise noted, this chapter is licensed under CC BY 4.0.

Chapter Information

ABSTRACT

This chapter presents marketing analytics as a form of decision support. It distinguishes decision support from reporting and from decision automation; describes the anatomy of a marketing decision, including alternatives, stakeholders, constraints, timing, and consequences; and introduces the analytic specification, a short working document that connects a marketing decision to analytical questions, data, candidate methods, deliverables, and a verification plan. Further sections examine how to match analytics types and deliverables to decisions, why baselines discipline interpretation, and why the cost of being wrong is rarely symmetric. A hands-on Google Colab lab turns a vague stakeholder request about StyleCraft's store expansion into a small, verified analysis. The chapter closes with managerial interpretation, industry practice, the ethics of problem framing, and exercises in specification writing and recommendation.

KEYWORDS

decision support; analytic specification; stakeholders; constraints; decision quality; baselines; marketing analytics; data-informed decision-making; artificial intelligence; verification

VERSION AND DATE

Version 1.0 · July 2026 · Language: English (United States)

SUGGESTED CITATION

Mendoza, J. (2026). Marketing analytics as decision support. In *Applied business analytics for marketing decision-making: Business analytics and data visualization* (Chapter 2, Version 1.0) [Open educational resource]. CC BY 4.0.

LICENSE AND RIGHTS

Copyright © 2026 Jose Mendoza. Except where otherwise noted, this work is licensed under the Creative Commons Attribution 4.0 International License (CC BY 4.0). You may share and adapt this material for any purpose, provided appropriate credit is given. Third-party trademarks, screenshots, figures, and other materials remain subject to their respective rights and licenses.

Google Colab is a product of Google LLC. “Python” and the Python logos are trademarks or registered trademarks of the Python Software Foundation. pandas is a sponsored project of NumFOCUS, a 501(c)(3) nonprofit charity in the United States. Tableau is a trademark of Salesforce, Inc. Product names are used for identification only and do not imply endorsement. StyleCraft Collective is a fictional company created for instruction.

COMPANION REPOSITORY

Datasets, notebooks, and figure sources for this chapter: <https://github.com/jrmst102/businessanalytics>

Chapter Learning Objectives

By the end of this chapter, students should be able to:

1. Explain what it means for analytics to function as decision support, and distinguish decision support from reporting and from decision automation.
2. Describe the anatomy of a marketing decision, including alternatives, stakeholders, constraints, timing, reversibility, and consequences.
3. Identify the stakeholders and constraints that shape an analytics project and explain why they must be identified before analysis begins.
4. Write an analytic specification that connects a marketing decision to analytical questions, a unit of analysis, data, candidate methods, a deliverable, and a verification plan.
5. Match the type of analytics and the form of the deliverable to the decision being supported.
6. Explain why analytic evidence requires a baseline and why the cost of being wrong is usually asymmetric.
7. Use a basic Google Colab workflow to turn a vague stakeholder request into a small, verified analysis and a qualified interpretation.

Chapter 1 introduced the role of the AI-augmented marketing analyst and the workflow that governs this guide: specify, predict-then-verify, explain, and document. This chapter examines the first and most consequential part of that workflow in depth. Before any data is loaded, any code is written, or any chart is drawn, the analyst must understand the decision the analysis is meant to support. The purpose of this chapter is to show how marketing analytics functions as decision support, and to introduce the working document that makes decision support disciplined: the analytic specification.

CONCEPT

What This Chapter Is Really About

Chapter 1 argued that analytics becomes valuable when it improves decisions. This chapter asks the practical question that follows from that claim: How does an analyst make sure that an analysis is actually connected to a decision? The answer developed here is that the connection must be designed, not assumed. The analyst identifies the decision, the decision-maker, the stakeholders, and the constraints, and then writes them down in an analytic specification before any code is run. In other words, the most important work in an analytics project often happens before the dataset is opened.

2.1 Marketing Decision Context: Three Requests and One Decision

Chapter 1 left the StyleCraft Collective marketing team at the end of a Monday morning meeting that had stalled. Everyone had data; not everyone had the same question. Two days later, the situation becomes more concrete. Three requests arrive in the analytics team's queue, all within the same afternoon.

The first request comes from the vice president of retail. It reads: "Finance keeps questioning the suburban and resort stores. Can you build us a dashboard showing that the new stores are working? We need it before the quarterly review."

The second request comes from the chief of staff to the chief marketing officer. It reads: "The CFO has asked whether we should pause the remaining store openings planned for next year. The CMO wants a recommendation by the 15th. What would it take to answer this properly?"

The third request comes from the CRM manager. It reads: "Quick one. Can you pull repeat purchase rates by store? No rush."

On the surface, these are three different tasks: a dashboard, a recommendation, and a data pull. Underneath, however, they appear to concern the same business issue: whether StyleCraft's expansion into affluent suburban and resort markets is succeeding. Yet the three requests frame that issue in very different ways. The first request presumes the answer; the analyst is asked to show that the stores are working, not to find out whether they are. The second request identifies a genuine decision, a decision-maker, and a deadline, but leaves the analytical work undefined. The third request specifies a metric but attaches it to no decision at all.

This situation is representative of professional practice. Analysts rarely receive well-formed analytical questions. They receive requests, and requests arrive shaped by the interests, assumptions, and vocabulary of the people who send them. The analyst's first task is therefore not technical. It is interpretive: to find the decision behind the request, and to negotiate a framing of the problem that the evidence can honestly serve.

In this chapter, these three requests will be used as a running example. The chapter will return to them after introducing the concepts needed to handle them well.

CONCEPT

Requests Versus Decisions

A request is what a stakeholder asks for: a dashboard, a report, a number, a chart. A decision is the choice the organization must make: continue, pause, expand, reallocate, redesign. Requests are the raw material of analytics work, but decisions are its purpose. When a request arrives without a decision attached, the analyst should ask what decision the output will inform. When a request arrives with the answer already attached, the analyst should ask whether the organization wants advocacy or evidence. This is an operational course concept rather than a standalone academic construct.

2.1.1 Opening Case Questions

1. What decision, if any, sits behind each of the three requests?
2. Which request would you handle first, and why?
3. The first request asks the analyst to show that the stores are working. What risks does this framing create for the analyst, the vice president, and the company?
4. What additional information would you need before agreeing to deliver the CMO's recommendation by the 15th?
5. The third request has no decision attached. Is it therefore harmless? What could go wrong if the analyst simply sends the numbers?

2.2 Reporting, Decision Support, and Decision Automation

The previous section presented three requests that concern the same underlying decision. This section introduces the vocabulary needed to classify what those requests are actually asking analytics to do. Analytics can relate to decisions in three broad ways: it can report on the business, it can support a human decision, or it can automate a decision. The three roles are related, but they place very different demands on the analyst.

Reporting describes the state of the business on a regular cadence. Dashboards, weekly campaign summaries, and monthly revenue reviews are reporting. Reporting is necessary, and much of it is descriptive analytics in the sense defined in Chapter 1 (Section 1.4). However, reporting is not automatically decision support. A dashboard that no decision depends on is monitoring, not support. Organizations can invest heavily in reporting without embedding the resulting insights into a specific decision process, which is one reason analytics investments often fail to change what the organization actually does (LaValle et al., 2011; Sharma et al., 2014).

Decision support connects evidence to a specific choice that a specific person or group must make. The analysis exists because a decision exists. The idea has a long lineage in information systems research, where decision support systems were defined from the start as systems that support, rather than replace, managerial judgment in semi-structured problems (Gorry & Scott Morton, 1971; Keen & Scott Morton, 1978). Davenport (2006) describes organizations that compete on analytics as those in which analysis is tied to distinctive capabilities and managerial action rather than produced for its own sake. In decision support, the human decision-maker remains responsible for the choice; the analysis narrows uncertainty, clarifies trade-offs, and makes the reasoning inspectable.

Decision automation removes the human from individual choices. A rule or model decides directly: a bidding algorithm sets the price of each ad impression, a send-time model chooses when each email is delivered, a recommendation engine selects which products appear on each customer's home page. Automation becomes a stronger candidate when decisions are repeated at scale, governed by a clear objective, supported by reliable data, rapidly measurable, and bounded by monitoring, guardrails, and a human override. Davenport et al. (2020) distinguish well-defined process automation from context-dependent managerial work and argue that AI is often

most valuable when it augments managers, while Huang and Rust (2021) locate repetitive and routine functions in the mechanical domain of AI. Deciding whether to pause a multi-year store expansion strategy has none of those properties. It is infrequent, high-stakes, slow to evaluate, and expensive to reverse. In other words, the StyleCraft expansion question is a decision-support problem, not an automation problem.

Table 2.1

Three Roles of Analytics in Marketing Organizations

Role	Question answered	Typical outputs	Marketing examples	Who decides
Reporting	How is the business doing?	Dashboards, scheduled reports, KPI reviews	Weekly campaign dashboard; monthly revenue by channel	May have no single decision attached; becomes decision support when tied to a defined review and response
Decision support	What should we do about a specific choice?	Analyses, memos, model results, experiment readouts	Whether to pause store expansion; how to reallocate holiday budget	A human decision-maker, informed by evidence
Decision automation	Which action should the system take in each case?	Rules, scores, and models embedded in operations	Real-time bidding; send-time optimization; product recommendations	A system, within limits set by humans

The boundary between the first two roles is not fixed. The same dashboard can be reporting for one audience and decision support for another. The difference is not the visual format; it is whether the output is attached to a choice, an owner, and a defined response. A weekly campaign report is ordinary monitoring when nothing follows from it, and decision support when it carries a review cadence, a threshold, an exception rule, and a named person who acts on it.

DEFINITION

Decision Support

Decision support is the use of data, analysis, models, and visualization to inform a specific managerial choice, while responsibility for the choice remains with a human decision-maker. In this guide, decision support is the default mode of marketing analytics: the analysis exists because a decision exists.

Source: Adapted from Gorry and Scott Morton (1971), Keen and Scott Morton (1978), Davenport (2006), and Sharma et al. (2014).

DEFINITION

Decision Automation

Decision automation is the delegation of individual, repeated decisions to rules, algorithms, or models, within boundaries and monitoring set by humans.

Course heuristic: marketing decisions are stronger candidates for automation when they are repeated at scale, governed by a clear objective, supported by reliable data, rapidly measurable, and bounded by monitoring, guardrails, and a human override. Low reversibility, serious consequences for individual customers, unstable objectives, or unclear accountability argue for human decision support rather than full automation.

Source: Course heuristic developed for this guide, informed by Davenport et al. (2020) and Huang and Rust (2021).

The distinction matters for the three StyleCraft requests. The dashboard request, taken literally, is a reporting task. The CMO's request is a decision-support task. The repeat-rate pull is a fragment that could serve either. The next section examines what a decision-support task requires the analyst to understand: the structure of the decision itself.

2.3 The Anatomy of a Marketing Decision

The previous section distinguished decision support from reporting and automation. This section examines what a decision actually consists of, because an analyst cannot support a decision without being able to describe it. Decision researchers have long argued that decisions are better understood as structured objects than as moments of choice: a decision has alternatives, values at stake, information, and consequences, and the quality of a decision can be assessed by how well those elements were handled (Keeney, 1992; March, 1994).

For marketing analytics, six elements are especially useful to identify.

First, the choice and its alternatives. A decision requires at least two genuine alternatives. "Should we pause the remaining store openings?" implies at least three: continue as planned, pause and reassess, or cancel. If only one alternative is acceptable to the organization, there is no decision to support, and the analyst should recognize the request as advocacy.

Second, the decision-maker. Someone must own the choice. For StyleCraft's expansion question, the CMO and CFO jointly own it; the vice president of retail is an interested party but not the decision-maker.

Third, the stakeholders. These are the people and groups affected by the decision or holding relevant information. Section 2.4 examines them in detail.

Fourth, the constraints. Time, budget, data availability, legal and privacy limits, and organizational realities all bound what analysis is possible. Constraints are examined in Section 2.4 as well.

Fifth, the timing and cadence. Some decisions are one-off and deadline-driven, such as the recommendation due on the 15th. Others recur on a cycle, such as monthly budget reallocation.

Recurring decisions justify reusable deliverables such as dashboards and models; one-off decisions usually justify a focused analysis and a memo. This distinction will matter again in Section 2.6.

Sixth, the consequences and reversibility. A steeply discounted flash sale can be evaluated within days and, if it fails, is unlikely to recur. A store lease commits capital for years. The less reversible and more consequential the decision, the more the analysis should invest in verification, in alternative explanations, and in honest statements of uncertainty.

Table 2.2
The Anatomy of a Marketing Decision, Applied to StyleCraft

Element	Question to ask	StyleCraft expansion example
Alternatives	What could the organization actually do?	Continue openings as planned; pause and reassess in two quarters; cancel remaining openings
Decision-maker	Who owns the choice?	CMO and CFO jointly, with board visibility
Stakeholders	Who is affected or holds information?	Retail, finance, CRM, brand, store staff, customers
Constraints	What bounds the analysis and the action?	Recommendation due by the 15th; store-level data only partially centralized; lease commitments already signed for two sites
Timing and cadence	Is this one-off or recurring?	One-off recommendation now; likely to recur as a quarterly review
Consequences and reversibility	How costly is being wrong, and can the action be undone?	Multi-year leases; brand positioning effects; pausing is partially reversible, canceling is not

CONCEPT

Decision Quality Versus Outcome Quality

A good decision can produce a bad outcome, and a bad decision can be rescued by luck. Decision quality refers to how well the choice was framed, informed, and reasoned at the time it was made; outcome quality refers to what later happened. Analysts support decision quality. This distinction protects analysts from two errors: claiming credit when a recommendation happens to work out, and abandoning sound reasoning because a defensible recommendation did not.

Source: Course concept developed for this guide, informed by Keeney (1992), March (1994), and Kahneman et al. (2011).

2.3.1 Levels of Marketing Decisions

It is also useful to classify marketing decisions by their level, a distinction that traces to Anthony's (1965) framework of strategic planning, management control, and operational control. Operational decisions are frequent and narrow: which subject line to send this week, which products to feature in Friday's drop. Tactical decisions allocate resources within an existing strategy: how to split the holiday budget across channels, which segments receive which offers. Strategic decisions change the strategy itself: whether to continue the suburban and resort expansion, whether to reposition the loyalty program. Moving up the levels, decisions become rarer, less reversible, and more dependent on judgment relative to data. Accordingly, the analyst's role shifts from producing repeatable outputs toward framing evidence, surfacing alternative explanations, and communicating uncertainty.

2.3.2 Good Enough, On Time

Finally, decision support is bounded by time. Simon (1997) observed that real decision-makers satisfice: they search for an option that is good enough given limits on time, information, and attention, rather than for a provably optimal one. The same logic applies to the analysis itself. An analysis that would be definitive in six weeks is useless for a recommendation due on the 15th. The professional skill is to design the best analysis that fits the deadline, to state clearly what the compressed analysis cannot establish, and to propose what should be measured next. In other words, "rigorous enough, on time, with stated limits" is usually more valuable than "perfect, late."

2.4 Stakeholders and Constraints

Section 2.3 identified stakeholders and constraints as two elements of every decision. This section examines each in more detail, because both shape the analysis before it begins.

2.4.1 Stakeholders

A stakeholder is any person or group that affects, or is affected by, the decision an analysis supports (Freeman, 1984). Stakeholders matter to the analyst for three practical reasons. First, they hold information: the retail lead knows that two of the newest stores opened with reduced marketing support, a fact that no dataset column records. Second, they hold criteria: finance evaluates the stores on contribution after occupancy costs, while brand evaluates them on positioning in affluent markets; the same evidence will be read differently against different criteria. Third, they hold influence: a technically sound recommendation that ignores a key stakeholder's concern is likely to be contested rather than used.

Furthermore, one stakeholder group is easy to overlook because it is not in the room: customers. Decisions about targeting, discounting, data use, and store closures affect customers directly, and their interests are not automatically represented by any internal function. Section 2.12 returns to this point.

Table 2.3

Common Stakeholders in a Marketing Analytics Project

Stakeholder	Typical criteria	What they contribute	What they may ask of the analyst
CMO / decision-maker	Growth, brand health, defensibility of the choice	The decision itself; strategic context	A clear, qualified recommendation by a deadline
Finance	Margin, contribution, payback	Cost data; scrutiny of assumptions	Evidence that revenue claims survive cost accounting
Retail operations	Store performance, staffing, leases	Operational context; store-level knowledge	Fair treatment of new stores still ramping up
CRM / lifecycle	Retention, repeat purchase, customer value	Customer-level history; program knowledge	Attention to long-term value, not only first purchase
Brand / creative	Positioning, customer perception	Qualitative context	Caution about metrics that ignore brand effects
Legal / privacy	Compliance, data governance	Boundaries on data use	Assurance that customer data is used appropriately
Data engineering	Data availability, definitions, pipelines	Access, documentation, metric definitions	Realistic timelines; consistent metric definitions

Stakeholder	Typical criteria	What they contribute	What they may ask of the analyst
Customers (silent stakeholders)	Value, trust, fair treatment	The behavior recorded in the data	Nothing directly; the analyst must represent them

2.4.2 Constraints

A constraint is any condition that bounds what the analysis or the decision can do. Constraints are sometimes treated as obstacles to be lamented. It is more productive to treat them as design inputs: they determine which analytical questions are answerable, which methods are feasible, and which deliverables are appropriate. An analyst who surfaces constraints at the start can design around them; an analyst who discovers them at the end has usually wasted the work.

Table 2.4
Common Constraints on Marketing Analytics Projects

Constraint type	Examples	Implication for the analysis
Time	Recommendation due by the 15th	Choose methods that fit the deadline; state what was not possible
Data	Store-level costs not centralized; short history for new stores	Some questions become unanswerable now; note them as next measurements
Budget and tools	No paid experimentation platform; free-tier Colab	Prefer analyses of existing data over new data collection
Skills	Team new to Python; AI assistance available	Keep methods explainable; verify AI-generated code
Legal and privacy	Customer data cannot leave approved systems	No pasting of customer records into external AI tools (see Section 1.13)
Organizational	The vice president has publicly championed the stores	Framing and communication require care; evidence must be evenhanded

CONCEPT

Constraints as Design Inputs

In this guide, constraints are treated as inputs to analytic design rather than as excuses. The analyst identifies time, data, budget, skill, legal, and organizational constraints before the work begins; chooses questions, methods, and deliverables that fit them; and records in the specification what the constraints made impossible. This record protects the analyst and improves the next iteration of the work.

So far, this chapter has described decisions, stakeholders, and constraints as things the analyst must understand. The next section introduces the document that captures that understanding in writing.

2.5 The Analytic Specification

The previous sections examined the parts of a decision. This section assembles them into the working document at the center of this chapter: the analytic specification. The specification is the written answer to the specify habit introduced in Section 1.7. Chapter 1 applied that habit to a single prompt; this chapter applies it to an entire project.

DEFINITION

Analytic Specification

An analytic specification is a short written document, typically one page, that records the decision to be supported, the decision-maker and deadline, the analytical questions, the unit of analysis, the data to be used, the candidate methods, the intended deliverable, the verification plan, the known limitations, and the planned use of AI assistance. It is written before the analysis begins and revised as the work proceeds. The specification is an operational concept for this guide, adapted from the business understanding phase of CRISP-DM.

Source: Adapted from Chapman et al. (2000) and Provost and Fawcett (2013).

The specification serves four purposes. First, it forces the translation from business question to analytical question, in the sense developed in Section 1.5, to happen explicitly and early. Second, it creates alignment: stakeholders can read one page and object to the framing before the work is done rather than after. Third, it disciplines AI use, because the specification supplies the context that turns a vague prompt into a strong one. Fourth, it creates accountability: when the analysis is finished, the results can be checked against the questions the specification promised to answer.

Table 2.5*The Analytic Specification Template, Applied to the StyleCraft Expansion Question*

Component	Question it answers	StyleCraft example (abridged)
Decision	What choice will this analysis inform?	Whether to continue, pause, or cancel the remaining store openings
Decision-maker and deadline	Who decides, and by when?	CMO and CFO; recommendation due the 15th
Analytical questions	What must be measured or compared?	How do stores compare on average order value, repeat purchase rate, and operating-cost-to-revenue ratio by store type? How do the newest stores compare with older stores at the same age?
Unit of analysis	What does one row represent?	One store (the most recent twelve months for stores open at least one year, and the full available period for stores open less than one year)
Data	What data, from where, covering what period?	Store summary extract from finance and CRM, covering the most recent twelve months for stores open at least one year, and the full available period for stores open less than one year
Candidate methods	What analysis fits the questions and the deadline?	Descriptive comparison by store type; cohort comparison by store age
Deliverable	What will the decision-maker receive?	A two-page memo with one summary table and a stated recommendation
Verification plan	How will the output be checked?	Row counts against the store roster; rate ranges between 0 and 1; totals reconciled to finance figures; one row hand-checked
Limitations and risks	What can this analysis not establish?	No causal claims; the reporting window is not equal across stores, so the newest store is not directly comparable; store costs exclude product cost

Component	Question it answers	StyleCraft example (abridged)
AI-use plan	Where will AI assist, and how will it be checked?	Draft pandas code and memo language; all code verified per Section 1.7; documented per Appendix D

Two features of the template deserve emphasis. The verification plan is written before the results exist, which is the predict-then-verify habit applied at project scale: the analyst commits in advance to what “checked” will mean. Moreover, the limitations component is written at the start, not retrofitted at the end. An analyst who knows from the outset that the data cannot support a causal claim is far less likely to draft one under deadline pressure.

2.5.1 Example: Rewriting the Dashboard Request

Consider the first request from Section 2.1: “Can you build us a dashboard showing that the new stores are working?” A specification-minded analyst would respond by proposing a reframing, in roughly the following form: “I understand the quarterly review is the deadline. To be useful there, I would propose an analysis of how the new stores are performing, on metrics finance will accept, with new stores compared against older stores at the same age. If the evidence is favorable, it will be far more persuasive coming from an evenhanded analysis than from a dashboard built to advocate.”

Notice what the reframing changes. The presumed conclusion (“the stores are working”) becomes an open question. The deliverable is negotiated from a standing dashboard, which suits a recurring decision, to a focused analysis, which suits the one-off review. Furthermore, the stakeholder’s legitimate need, being well prepared for the quarterly review, is preserved. Reframing is not refusal; it is the conversion of a request into a decision the evidence can honestly serve.

2.6 Matching Analytics and Deliverables to the Decision

The previous section introduced the specification, which includes a candidate method and a deliverable. This section examines how those two choices should be made. Chapter 1 (Section 1.4) introduced the four types of analytics: descriptive, diagnostic, predictive, and prescriptive. The specification asks the analyst to choose among them deliberately, based on the decision rather than on ambition.

A useful discipline is to begin with the least sophisticated analysis that could inform the decision, and to escalate only when the decision requires it. Many managerial decisions are adequately supported by careful descriptive comparison. The StyleCraft expansion question, for instance, begins descriptively: how do the store types actually compare on the metrics finance cares about? Diagnostic work follows if the comparison raises questions of why. Predictive work is justified when the decision depends on what will happen, such as forecasting whether new stores will reach contribution targets. Prescriptive decisions may require causal evidence from

experiments when the organization must choose among actions and the cost of choosing incorrectly is high; Chapter 11 examines experiments and causal evidence directly. It is worth separating the two ideas here. An experiment estimates the causal effect of the alternatives; it does not by itself prescribe the action. The prescription comes from combining the experimental result with objectives, costs, constraints, and judgment. In other words, an experiment is a way of generating evidence, not a fifth type of analytics.

The deliverable should likewise follow the decision, and especially its cadence. One-off decisions are usually best served by a memo or briefing: a short document with the recommendation, the evidence, the limitations, and the next measurement. Recurring decisions justify reusable deliverables: a dashboard for monitoring, a model score refreshed on a schedule, an experiment pipeline. Building a dashboard for a one-off decision wastes effort; writing a one-off memo for a decision that recurs monthly wastes the organization's time twelve times a year.

Table 2.6
Matching the Decision to the Analytics and the Deliverable

Decision situation	Evidence or analytics needed first	Typical deliverable	Cadence
Quarterly review of store performance	Descriptive comparison	Analysis memo with summary table	One-off, likely to recur quarterly
Understanding why repeat purchase declined	Diagnostic (cohorts, segments)	Analysis memo or briefing	One-off
Monitoring campaign health	Descriptive	Dashboard	Recurring (weekly)
Deciding which customers receive a retention offer	Predictive (propensity)	Model score plus targeting rule	Recurring (each campaign)
Choosing between two offer designs	Causal evidence from a randomized experiment	Experiment readout with lift, uncertainty, and a decision recommendation	One-off per test
Setting the holiday budget split	Prescriptive (allocation with judgment)	Recommendation memo with scenarios	Recurring (annual)

CONCEPT

Decision Cadence

Decision cadence is the frequency with which a decision recurs. Cadence drives the choice of deliverable: one-off decisions favor focused analyses and memos, while recurring decisions justify reusable assets such as dashboards, refreshed model scores, and standing experiments. Mismatched cadence is a common and expensive analytics failure. This is an operational course concept rather than a standalone academic construct.

2.7 Evidence, Baselines, and the Cost of Being Wrong

Section 2.6 matched analyses to decisions. This section examines two properties that any decision-supporting evidence must have: a baseline, and an honest accounting of what being wrong would cost.

2.7.1 Compared to What?

A number by itself supports no decision. Suppose the analyst reports that StyleCraft's suburban stores have an average order value near \$119. Is that good? The question is unanswerable until a baseline is chosen: compared to the urban stores? To the same stores last year? To the plan finance approved? To what those customers would have spent online anyway? Each baseline answers a different analytical question, and some baselines flatter while others do not. The suburban stores look strong against an average-order-value baseline and weak against a repeat-purchase baseline, which is precisely why the Chapter 1 meeting stalled.

CONCEPT

Baseline

A baseline is the comparison point against which a metric or result is judged: a prior period, a comparable group, a plan or target, a do-nothing alternative, or a stated expectation. Every claim of the form "X is high," "X improved," or "X is working" implicitly depends on a baseline, and the analyst should make that baseline explicit. Choosing the baseline before seeing the results is a form of predict-then-verify.

Source: Course concept developed for this guide, informed by Provost and Fawcett (2013) and Shmueli (2010).

Baselines also discipline the treatment of new stores. Comparing a store seven months after opening against a store in its sixth year answers almost nothing; comparing it against older stores at the same age is a fairer analytical question. The lab in Section 2.9 makes this concrete.

2.7.2 Asymmetric Errors

Every recommendation can be wrong in two directions, and the two directions rarely cost the same. If StyleCraft pauses an expansion that would have succeeded, it loses growth, momentum

with landlords, and credibility for the strategy. If it continues an expansion that is failing, it compounds multi-year lease losses. The analysis cannot eliminate either error, but the recommendation should acknowledge which error is more costly and more reversible, because that asymmetry legitimately influences the decision even when the evidence is ambiguous (Kahneman et al., 2011; March, 1994).

The same logic recurs throughout marketing analytics at smaller scale. A retention model that wrongly flags a loyal customer as at-risk costs one unnecessary discount; a model that misses a genuinely departing customer costs the entire future relationship. Chapter 9 formalizes this reasoning as false positives and false negatives; the managerial habit of asking “which mistake hurts more?” is available now, long before the formal machinery.

2.7.3 Is the Analysis Worth Doing?

Finally, evidence earns its cost by improving the decision it informs. Before commissioning the analysis, the analyst should ask whether any plausible result could change the action, its timing, its scale, its safeguards, or the confidence with which it is taken. If no result could affect any of those, the analysis has little decision value. If the store leases for next year are already signed and cannot be exited, an elaborate analysis of whether to sign them has no decision value, however interesting it may be. Conversely, if a modest analysis could move the choice between pause and continue, it is worth doing well. Formally, the value of information is the improvement in expected decision payoff obtained by reducing uncertainty, net of the cost of obtaining that information (Howard, 1966; Keeney, 1992). Practically, the question is the professional defense against producing analytics for its own sake (Pfeffer & Sutton, 2006).

2.8 AI as a Decision-Support Assistant

The previous sections have been about framing. This section asks where AI tools fit into framing work, because the four habits from Section 1.7 apply here just as they apply to code.

AI assistants can be genuinely useful before any data is touched. Given a description of the situation, an assistant can draft a first version of the analytic specification, propose analytical questions the analyst had not considered, enumerate plausible stakeholders and constraints, and suggest baselines. Used this way, the assistant functions as a brainstorming partner that widens the analyst’s view of the problem.

However, the risks at the framing stage are distinctive. A drafted specification can be fluent, plausible, and generic: it may propose questions that fit any retailer rather than StyleCraft specifically, or omit the organizational constraint that no dataset records. More subtly, an assistant will usually accept the framing it is given. Asked to “outline an analysis showing the new stores are working,” most assistants will comply, producing a professional-looking plan for advocacy. The assistant amplifies the framing; it does not audit it. Because AI outputs are sensitive to the prompt and may produce fluent but generic plans, the analyst should treat AI-generated specifications as drafts to be audited, not as authority. Responsibility for the framing therefore remains entirely with the analyst.

A practical countermeasure is to use the assistant against the framing deliberately. After drafting a specification, ask the assistant to argue the other side: What would a skeptical CFO say about this plan? Which analytical questions are missing? What alternative explanations could produce the same evidence? This use of AI, sometimes described as red-teaming one's own work, converts the tool from an accelerator of the first idea into a test of it.

AI IN PRACTICE

Ask for the Other Side

Before finalizing an analytic specification, submit it to an AI assistant with a prompt of the form: "Here is my analysis plan. Act as a skeptical reviewer. Identify the weakest analytical question, one stakeholder I have ignored, one alternative explanation for the result I expect, and one verification check I have not planned." The assistant's answer is not authoritative, but it is cheap, fast, and frequently catches a genuine gap. Document the exchange as described in Appendix D.

2.9 Hands-On Application in Python and Google Colab

The preceding sections were conceptual. This section moves from concepts to practice by carrying one of the StyleCraft requests through a small analysis. As in Chapter 1, the lab uses a compact dataset created directly in Python so that the focus remains on the workflow rather than on file handling. The store figures below are a simplified preview of StyleCraft's full store and transaction data; later chapters load the complete, realistic datasets from the companion repository. The numbers are constructed for instruction and should not be treated as industry benchmarks.

2.9.1 Lab 2.1: From Request to Specification to First Evidence

The scenario is the CMO's request from Section 2.1: a recommendation on the store expansion is due by the 15th. Assume the specification in Table 2.5 has been agreed upon. The first analytical question reads: How do stores compare on average order value, repeat purchase rate, and operating-cost-to-revenue ratio by store type? The unit of analysis is one store, summarized over the most recent twelve months for stores open at least one year, and the full available period for stores open less than one year. That definition is deliberate rather than tidy. Scottsdale has been open seven months, so its row covers seven months and not twelve, and its totals are reported as observed rather than annualized. The windows are therefore not equal across stores, which is a comparability limitation the specification records before any number is produced.

Before writing any code, apply predict-then-verify. The expectations below are written down first, so that the output has something to be checked against.

VERIFICATION CHECK

Before you run: the table should contain seven stores and eight columns. StyleCraft's catalog runs from roughly \$24 to \$90 a unit and baskets contain one or more items, so average order value should be positive and should plausibly fall below a few hundred dollars. Repeat purchase rate, expressed as a proportion, must fall between 0 and 1. Based on the Chapter 1 meeting, expect the suburban and resort stores to show higher average order values but lower repeat rates than the urban stores.

After you run: the frame should report seven rows and eight columns, no column should contain missing values, and every store should appear exactly once. The metric columns computed in Code 2.3 should then fall inside the ranges predicted above.

Investigate if: a repeat rate falls outside 0 and 1, an average order value falls outside the predicted range, or an operating-cost ratio exceeds 1. Note that the catalog floor is not a floor for the metric, because discounts, returns, credits, and partial refunds can pull a realized average order value below the cheapest list price. A ratio above 1 is not an arithmetic error, but it must be explained before the table is discussed.

Code 2.1. Create the store summary table

```
import pandas as pd

stores = pd.DataFrame({
    "store": ["SoHo", "West Village", "Brooklyn", "Rye",
             "Greenwich", "Hamptons", "Scottsdale"],
    "store_type": ["Urban", "Urban", "Urban", "Suburban",
                  "Suburban", "Resort", "Resort"],
    "months_open": [60, 48, 40, 30, 24, 20, 7],
    "revenue": [3240000, 2150000, 1890000, 1610000,
               1780000, 1420000, 610000],
    "orders": [49850, 33600, 30050, 13650, 14850, 11450, 5050],
    "customers": [20600, 14200, 13900, 9150, 10300, 8900, 4300],
    "repeat_customers": [9680, 6320, 5890, 1690, 1850, 1160, 310],
    "operating_cost": [2050000, 1430000, 1260000, 1330000,
                      1490000, 1260000, 720000]
})

stores
```

Expected output: a table with seven rows, one per store, and eight columns.

Each row represents one store. The columns record the store type, the number of months the store has been open, revenue for the store's reporting window, order and customer counts, the number of customers who purchased more than once in the window, and the store's operating cost. One definition deserves attention before any metric is computed: `operating_cost` here includes occupancy, staffing, and local marketing, but excludes the cost of the products sold; the name is chosen deliberately so that no reader mistakes it for total cost. This matters for interpretation later, and noticing it now is part of the specify habit.

2.9.2 Verifying Structure Before Computing Metrics

The first verification checks concern structure, not conclusions.

Code 2.2. Run basic structure checks

```
# Expect 7 rows and 8 columns
stores.shape

# Expect no missing values
stores.isna().sum()

# Expect every store in the roster exactly once
stores["store"].is_unique
```

Expected output: (7, 8); a column of zeros; True.

If the shape, missingness, or uniqueness checks fail, the analysis stops until the discrepancy is explained. In a realistic project, this is also the point at which totals would be reconciled against finance's own figures, as the verification plan in Table 2.5 promised.

2.9.3 Computing the Decision Metrics

The next cell computes the three metrics named in the analytical question, plus revenue per customer as supporting context.

Code 2.3. Calculate store-level decision metrics

```
stores["aov"] = stores["revenue"] / stores["orders"]
stores["repeat_rate"] = stores["repeat_customers"] / stores["customers"]
stores["operating_cost_ratio"] = stores["operating_cost"] / stores["revenue"]
stores["revenue_per_customer"] = stores["revenue"] / stores["customers"]

stores[["store", "store_type", "months_open",
        "aov", "repeat_rate", "operating_cost_ratio"]].round(3)
```

Expected output (selected rows): SoHo shows aov 64.995, repeat_rate 0.470, operating_cost_ratio 0.633; Rye shows 117.949, 0.185, 0.826; Scottsdale shows 120.792, 0.072, 1.180.

Now compare the output with the predictions. Average order values run from roughly \$63 to \$65 for the urban stores and from roughly \$118 to \$124 for the suburban and resort stores, all positive and all well below the predicted ceiling. Repeat rates run from about 0.42 to 0.47 in urban stores, about 0.18 in suburban stores, and lower still in the resort stores, all between 0 and 1. The direction matches the expectation from the Chapter 1 meeting. One value stands out: Scottsdale's operating-cost ratio exceeds 1, meaning the store's operating cost exceeded its revenue in the window. A ratio above 1 is not an arithmetic error, but it is exactly the kind of value that must be flagged, because it will dominate any discussion of the table.

2.9.4 Summarizing by Store Type, Carefully

The analytical question asks for comparison by store type. There are two ways to aggregate, and they answer different questions, a distinction first raised in Exercise 1.10, item 2, “The Averaged Average.”

Code 2.4. Compare equal-weighted and pooled aggregation

```
# Equal-weighted average of store-level rates (each store counts equally)
stores.groupby("store_type")[["aov", "repeat_rate",
                             "operating_cost_ratio"]].mean().round(3)

# Weighted (pooled) rates built from the underlying totals
totals = stores.groupby("store_type")[["revenue", "orders",
                                       "customers", "repeat_customers",
                                       "operating_cost"]].sum()
totals["aov"] = totals["revenue"] / totals["orders"]
totals["repeat_rate"] = totals["repeat_customers"] / totals["customers"]
totals["operating_cost_ratio"] = totals["operating_cost"] / totals["revenue"]
totals[["aov", "repeat_rate", "operating_cost_ratio"]].round(3)
```

Expected output, pooled version: Urban 64.141, 0.449, 0.651; Suburban 118.947, 0.182, 0.832; Resort 123.030, 0.111, 0.975. The equal-weighted averages differ slightly, for example Resort 122.405, 0.101, 1.034.

The equal-weighted mean answers a per-store question: what does the average store of this type look like when every store receives the same weight? The pooled version answers a portfolio question: what is the aggregate rate for this store type when each metric is reconstructed from its underlying totals? Note that each pooled metric is weighted by its own denominator: average order value by orders, repeat rate by customers, and the operating-cost ratio by revenue. The habit to build is therefore not only “which denominator?” but “which unit receives the weight?” — the store, the order, the customer, or the revenue dollar. Chapter 3 returns to that question when it defines grain and the elements of a metric formally. With these seven stores the two versions differ only modestly, for example a resort repeat rate of about 0.10 under the equal-weighted mean versus about 0.11 pooled, but the gap widens whenever group members differ greatly in size. For a finance audience evaluating the economics of the expansion, the pooled version is usually the right default.

The pooled results tell a compact story. Urban stores show an average order value near \$64, a repeat rate near 0.45, and an operating-cost ratio near 0.65. Suburban stores show an average order value near \$119, a repeat rate near 0.18, and an operating-cost ratio near 0.83. Resort stores show the highest average order value, near \$123, the weakest repeat rate, near 0.11, and an operating-cost ratio near 0.98, meaning operating costs consume nearly all revenue before the cost of the products themselves is even counted. Section 2.10 asks what may, and may not, be concluded from this.

2.9.5 One More Verification: Hand-Check a Row

The final planned check is a hand computation. For Rye: 1,610,000 divided by 13,650 orders is approximately \$117.9 of average order value; 1,690 repeat customers divided by 9,150 customers is approximately 0.185; 1,330,000 divided by 1,610,000 is approximately 0.826. All three match the computed table. A single hand-checked row does not prove the table correct, but it catches reversed numerators, wrong columns, and unit errors, which are among the most common failures in practice, and it costs one minute.

2.10 Marketing Interpretation and Managerial Insight

The previous section produced a verified table. This section asks what the table means for the decision, which is where decision support succeeds or fails.

A weak interpretation would say: “The new stores are losing money and customers do not come back. StyleCraft should cancel the expansion.” This reading overreaches in at least three ways. First, it treats the resort operating-cost ratios as a settled verdict, when Scottsdale has been open seven months: a store that new has had little time to accumulate repeat customers, and its costs include one-time opening expenses. Its repeat rate of about 0.07 is partly a mechanical consequence of its age, which is why the specification’s second analytical question, comparing new stores against older stores at the same age, exists. Second, the table says nothing about why suburban customers do not return: the data cannot distinguish a customer who defected from one who shifted to the online channel, a distinction the current unit of analysis, the store, cannot see. Third, no causal claim is available: the stores differ in age, location, customer mix, and marketing support, so the comparison describes differences without explaining them.

A stronger interpretation would say: “Across each store’s reporting window, the suburban and resort stores generate substantially higher average order values than the urban stores, roughly \$119 to \$123 versus \$64, but convert far fewer customers into repeat buyers, roughly 0.11 to 0.18 versus 0.45, and operate at much heavier operating-cost ratios, roughly 0.83 to 0.98 versus 0.65, before product costs. The newest store’s figures are not yet comparable because of its short history. The evidence is consistent with the concern that the expansion’s economics depend on repeat behavior that has not materialized, but it does not establish why, and it does not yet compare stores at the same age.”

The stronger interpretation is better because it is specific, it is bounded by the unit of analysis, it separates what the evidence shows from what it suggests, and it leaves the decision with the decision-maker while making the decision easier to reason about.

2.10.1 Revisiting the StyleCraft Case

Return, finally, to the three requests from Section 2.1. The CMO’s request is served by the memo the specification promised: the pooled comparison, the qualified interpretation above, a recommendation, and two named next measurements, a same-age store comparison and a customer-level analysis of whether suburban buyers repurchase online. The vice president’s

request is served by the reframing in Section 2.5.1; the evenhanded analysis is what makes any favorable evidence credible at the quarterly review. Moreover, the CRM manager's request can now be answered responsibly: the repeat rates can be sent along with the store-age caveat attached, so that a number pulled "with no rush" does not circulate stripped of the context that makes it interpretable. In other words, one specification has organized three requests around one decision.

The recommendation itself might reasonably be: pause the two most discretionary openings pending the same-age comparison, continue the committed sites, and decide again next quarter with the two named analyses in hand. Whether the CMO accepts that recommendation depends on judgment, risk appetite, and considerations beyond the table, which is exactly how decision support is supposed to work.

2.11 Business Analytics in Practice

The previous section interpreted a small table for a fictional retailer. This section leaves StyleCraft behind and looks at how the framing habits developed in this chapter appear in real organizations: where the specification actually lives in a company's weekly rhythm, what the right deliverable is when the deadline is measured in hours, and where the boundary between decision support and decision automation falls in practice.

2.11.1 Where the Specification Lives

Most analytics requests do not arrive by email. In retail, a great many of them arrive in the weekly trading meeting, a standing review in which merchandising, planning, marketing, e-commerce, and finance work through the trading week together: what sold, what did not, what is on markdown, what is short of stock, and what the coming week's plan should be. The meeting exists because retail decisions are made on a weekly cadence, and it functions in practice as a decision queue. Someone observes that a category is behind plan, someone else offers an explanation, and within a few minutes an analyst has been asked for something.

What separates teams that use analytics well from teams that merely produce it is what happens in those few minutes. On weaker teams the request leaves the room as a deliverable: send me the outerwear numbers. On stronger teams it leaves as a specification, even an informal one. The analyst asks who owns the decision, what will be decided, and by when. Is the meeting deciding whether to mark down, whether to reorder, or whether to move media budget? Each of those is a different analytical question, with a different deadline and a different acceptable deliverable, and the outerwear numbers serve at most one of them well.

This is not an abstract preference. LaValle et al. (2011) found that the organizations obtaining value from analytics were those that began with business questions rather than with available data, and that embedded the resulting insights into an actual business process. Sharma et al. (2014) argue similarly that the returns to analytics depend on the organizational decision-making processes surrounding the analysis rather than on the analytical technology itself. A weekly trading meeting is exactly such a process. The specification does not have to be a formal

document to do its work; it has to exist before the analyst leaves the room. The practice lesson is that in most organizations there is a recurring meeting where decisions are actually made, and locating it is among the highest-return things a new analyst can do.

2.11.2 Forty-Eight Hours, and the Right Answer Was a Memo

Not every decision leaves room for the analysis an analyst would prefer. A common situation in marketing is the mid-quarter reallocation: a channel is underdelivering, a competitor has moved, or a budget line must be committed before the end of the week, and the analytics team has two days. Two days is not enough to build a marketing mix model or to run an incrementality test. It is enough to assemble what is already known, state the assumption behind each number, and write down what the evidence can and cannot support.

Under that constraint the strongest deliverable is usually a memo rather than a model. That is a less modest claim than it sounds. Amazon, a company with no shortage of modeling capacity, does not use slide decks for decisions of consequence. “We don’t do PowerPoint (or any other slide-oriented) presentations at Amazon,” Bezos wrote in the 2017 letter to shareholders. “Instead, we write narratively structured six-page memos. We silently read one at the beginning of each meeting in a kind of study hall” (Bezos, 2018). The reasoning is that prose forces the writer to make the argument explicit, including the parts that a bulleted chart can leave comfortably ambiguous.

For a forty-eight-hour reallocation request, a memo does several things a rushed model cannot. It states the decision and the alternatives. It presents the evidence with its baseline attached. It says plainly which comparisons are confounded and which figures are estimates. And it names the measurement that would settle the question given more time, which converts an admitted limitation into a proposal. Simon’s (1997) account of satisficing applies directly: the decision-maker is not searching for the optimal answer but for the best defensible one available before the money is committed. The practice lesson is that matching the deliverable to the deadline is an analytic decision rather than an administrative one. A model delivered after the budget is committed is worth nothing, however good it is.

2.11.3 The Automation Boundary

Section 2.2 distinguished decision support from decision automation. In practice that boundary is drawn far more sharply than most conversations about AI suggest, and marketing offers an unusually clean pair of examples.

At one end is search advertising. Google’s automated bidding is explicit about what it does: it sets bids for each individual auction rather than a few times a day, using signals available at auction time, including device, location, time of day, audience list, and query text, together with combinations of those signals (Google, n.d.). No human could participate at that frequency. The objective is stated numerically, each auction is individually inexpensive, feedback arrives within hours, and a bad setting can be changed immediately. Every condition in this chapter’s automation heuristic is satisfied. The analyst’s job here is not to set bids. It is to define the

conversion the system optimizes toward, to watch the guardrails, and to notice when the objective the system is maximizing has stopped matching the objective the business cares about.

At the other end is brand repositioning. In January 2009, Tropicana relaunched its Pure Premium line with a redesigned carton that replaced the familiar orange-and-straw image. Sales of the line fell roughly 20 percent against the same period a year earlier, the company reverted to the previous packaging by the end of February, and a subsequent demand study estimated the cost of the episode at approximately \$27 million (Lee et al., 2010). That decision was infrequent, expensive, slow to evaluate, and, once shelf presence and customer perception had moved, only partly reversible. It also depended on something no dataset contained: how much of the brand's value sat in a picture customers used to find the product.

Nothing in the automation heuristic would have permitted delegating that choice to a model, and nothing in the bidding case would have justified deciding each auction by hand. The two examples differ not in analytical sophistication but in the structure of the decision. Davenport et al. (2020) draw the same line between well-defined process automation and context-dependent managerial work, and Huang and Rust (2021) locate repetitive, routine functions in the mechanical domain of AI while reserving judgment-dependent work for people. The practice lesson is that the question is never whether AI is capable. It is whether the decision has the structure that makes delegation responsible.

2.11.4 In Your First Analyst Job

Across these three situations, one failure recurs more often than any technical error: answering the wrong question well. It is a comfortable failure, because everything about the work looks correct. The code runs, the chart is clean, the totals reconcile, and the analyst can defend every step. The only defect is that the output does not bear on the choice the organization is making, and that defect is close to invisible from inside the analysis.

The protection is unglamorous. Before beginning, find the decision, the person who owns it, and the date by which it must be made, then write those three things down where the stakeholder can see them and disagree. Ask what the deliverable will be used for, not only what it should contain. When the deadline is too short for the analysis you would prefer, say what you can support and what you cannot, and propose the measurement that would close the gap.

In your first analyst job, no one will hand you the specification. Requests will arrive as requests, in meetings, in messages, and in hallways, and they will sound urgent. The habit of converting them into decisions before opening the data is what separates an analyst from a person who produces output.

2.12 Ethics, Framing, and Responsible Decision Support

So far this chapter has treated framing as a matter of effectiveness. This section treats it as a matter of responsibility, extending the discussion of responsible analytics from Section 1.13.

The first issue is predetermined framing. The request to “show that the stores are working” asks the analyst to begin from the conclusion. Producing advocacy dressed as analysis misleads

the decision-maker, exposes the organization to a worse decision, and spends the analyst's credibility, which is the professional asset everything else depends on. The responsible response, illustrated in Section 2.5.1, is to reframe toward an open question while honoring the stakeholder's legitimate need.

The second issue is metric gaming. When a measure becomes a target, behavior reorganizes around the measure and the measure degrades as evidence (Campbell, 1979; Strathern, 1997). If store managers learn that repeat rate decides the expansion's fate, enrollment tactics that inflate measured repeat purchases without creating real loyalty become tempting. Analysts should anticipate this: prefer a small set of complementary metrics over a single target, and treat sudden improvement in a targeted metric as a finding to investigate rather than a success to announce.

The third issue is who the framing serves. Every stakeholder in Table 2.3 is represented in the room except the customers whose behavior constitutes the data. Decisions framed purely on contribution can quietly disadvantage groups of customers, for example by withdrawing stores, service, or offers from the markets where the company has historically invested least, an echo of the disparate-impact concern raised in Section 1.13 (Barocas & Selbst, 2016). Privacy obligations likewise do not appear in any metric, yet they bind the analysis throughout (Martin & Murphy, 2017). A responsible specification names these considerations explicitly, so that they are weighed rather than omitted.

CONCEPT

Responsible Framing

Responsible framing is the practice of posing analytical questions as open questions rather than predetermined conclusions, choosing metrics with awareness of how they may be gamed, and representing affected parties, including customers, who are not present when the decision is framed. It extends responsible analytics (Section 1.13) from the use of data to the design of the question.

Source: Course concept developed for this guide, informed by Campbell (1979), Strathern (1997), Barocas and Selbst (2016), and Martin and Murphy (2017).

2.13 Chapter Summary

This chapter developed the claim, introduced in Chapter 1, that analytics becomes valuable when it improves decisions. The main point is that the connection between analysis and decision must be designed. Analytics can report, support, or automate; this guide's default mode is decision support, in which evidence informs a choice that remains with a human decision-maker. A decision has an anatomy, alternatives, a decision-maker, stakeholders, constraints, a cadence, and consequences, and the analyst must be able to describe it before supporting it.

The chapter's central instrument is the analytic specification: one page, written before the work begins, recording the decision, the analytical questions, the unit of analysis, the data, the candidate methods, the deliverable, the verification plan, the limitations, and the AI-use plan. The chapter also argued that every claim needs an explicit baseline, that the two ways of being

wrong rarely cost the same, and that an analysis is worth doing only if some result could change the decision. The lab practiced the full arc on StyleCraft’s store expansion: request, specification, verified evidence, qualified interpretation, and recommendation. The chapter then stepped outside StyleCraft to see where the specification lives in a real weekly trading meeting, why a forty-eight-hour budget request is often best answered with a memo rather than a model, and where the boundary between automated bidding and brand repositioning falls.

Looking ahead, the specification requires the analyst to name a unit of analysis, variables, and metrics, and to trust the measurements behind them. The next chapter examines exactly those foundations: data, measurement, and marketing variables, including metrics, KPIs, dimensions, features, validity, reliability, and the data dictionary that makes a dataset usable by a team.

2.14 Exercises for Practice and Homework

The following exercises practice the chapter’s main habits: find the decision behind the request, describe its anatomy, identify stakeholders and constraints, write the specification, choose the baseline, and connect verified evidence to a qualified recommendation. They are organized into three groups. Core chapter practice is the required path and should be completed by every student. In-class activities are designed for discussion. Extensions are optional and go beyond the required material. Each exercise also carries its assignment label (required practice, in-class discussion, homework submission, or optional) so that instructors can assign selectively.

2.14.1 Core Chapter Practice

Exercise 2.1 Concept Check (Required Practice)

1. Explain the difference among reporting, decision support, and decision automation, with one marketing example of each.
2. What is the difference between a request and a decision? Why does the difference matter?
3. List the components of an analytic specification and explain why the verification plan is written before the results exist.
4. Why does every claim of the form “the metric improved” depend on a baseline?
5. What is the difference between decision quality and outcome quality, and why does the distinction protect analysts?
6. Why is deciding whether to pause a store expansion a poor candidate for decision automation?

Exercise 2.2 Classify the Decision (Required Practice)

For each situation, classify the decision as operational, tactical, or strategic; state whether it is one-off or recurring; and name the deliverable you would propose.

1. Choosing the subject line for this week’s new-arrivals email.

2. Setting the split of the holiday budget across paid social, paid search, and influencer spend.
3. Deciding whether to reposition the loyalty program around experiences rather than discounts.
4. Selecting which lapsed customers receive next month's win-back offer.
5. Monitoring whether campaign conversion rates are drifting downward.
6. Deciding whether to open a second resort store in Florida.
7. Choosing which products appear in Friday's drop.
8. Reviewing store performance for the quarterly business review.

Exercise 2.3 Identify Stakeholders and Constraints (Required Practice)

For each situation, name at least three stakeholders, the criteria each would apply, and two constraints the analyst should surface before starting.

1. The CMO asks whether StyleCraft should shift budget from influencer marketing to paid search.
2. Finance asks whether the loyalty program's discounts are eroding margin.
3. The retail lead asks for an analysis of staffing levels against store traffic.
4. Legal asks whether a proposed lookalike-audience campaign uses customer data appropriately.
5. The CRM manager asks which customers should be excluded from discount campaigns.

Exercise 2.4 Write the Analytic Specification (Homework Submission)

Choose the second or third request from Section 2.1, or a request supplied by your instructor. Using the template in Table 2.5, write a one-page analytic specification. Every component must be completed, including the verification plan, the limitations, and the AI-use plan. Then write one paragraph explaining which component was hardest to complete and why.

Exercise 2.5 Choose the Baseline (Required Practice)

For each claim, propose two different baselines, state how each baseline changes the analytical question, and identify which baseline a skeptical CFO would prefer.

1. "The Greenwich store had a strong year."
2. "Repeat purchase improved after the loyalty redesign."
3. "Our TikTok acquisition costs are low."
4. "The Scottsdale store is underperforming."
5. "Email is our best channel."

Exercise 2.6 Hands-On Colab Practice (Homework Submission)

1. Open a new Google Colab notebook and reproduce Lab 2.1.
2. Add a column called `revenue_less_operating_cost` that subtracts `operating_cost` from `revenue`. Because product cost, corporate overhead, and taxes are excluded, this figure is

a partial operating comparison rather than contribution or profit, and the column name should say so. Before running it, predict which store type will rank last.

3. Recompute the pooled store-type summary including revenue_less_operating_cost. Does the ranking match your prediction?
4. Compute the simple and pooled repeat rates for the resort stores and report both. Explain in two sentences why they differ and which one you would show a finance audience.
5. Remove Scottsdale from the table and recompute the resort summary. Write two sentences on how much the newest store was driving the resort figures, and what that implies for the same-age comparison proposed in the specification.
6. Write one verification check you performed and one limitation of this dataset for the expansion decision.

Exercise 2.7 AI-Assisted Practice (Homework Submission)

1. Submit your Exercise 2.4 specification to an AI assistant using the red-teaming prompt from Section 2.8. Record what the assistant identified.
2. Revise the specification in response to at most two of the assistant's points, and reject at least one point with a one-sentence justification.
3. Document the exchange using the full AI-use documentation template in Appendix D.

Exercise 2.8 Managerial Memo Exercise (Homework Submission)

Write a memo of 250–350 words from the analytics team to the StyleCraft CMO and CFO, dated before the 15th, recommending a course of action on the remaining store openings. Use the Lab 2.1 evidence. The memo must state the decision, the recommendation, the evidence with its baseline made explicit, one limitation, the asymmetry between the two ways of being wrong, and two named next measurements. Recommended structure: decision context, evidence, recommendation, limitations and risks, next steps.

2.14.2 In-Class Activities

Exercise 2.9 Find the Flaw in the Framing (In-Class Discussion)

1. The missing decision. A team spends three weeks building a beautiful customer dashboard. At the demo, an executive asks, “So what should we do differently?” and no one can answer. Where in this chapter's terms did the project go wrong?
2. The convenient baseline. An analysis reports that the resort stores' revenue grew 40 percent year over year, without mentioning that the prior year included only a partial season after opening. Which baseline was chosen, and which was avoided?
3. The single metric. Leadership announces that store bonuses will now depend on repeat purchase rate. Predict two behaviors this may produce, and explain the risk in this chapter's terms.

4. The wrong cadence. An analyst answers the monthly budget-allocation question with a bespoke forty-page analysis every month. What is the mismatch, and what should be built instead?
5. The unanswerable question. The CMO asks whether the Hamptons store “would have succeeded” if it had opened two years earlier. What can and cannot be said, and what would you offer instead?
6. The overreached memo. A draft memo states: “The data proves that suburban customers are disloyal.” Rewrite the sentence so that it stays within what the Lab 2.1 evidence supports.

Exercise 2.10 Ethics and Framing Mini-Cases (In-Class Discussion)

For each case, identify the issue, the stakeholders affected, and what a responsible analyst should do instead.

1. The friendly deadline. A stakeholder offers to extend the analyst’s deadline by a week “if the numbers can come out looking a bit stronger.”
2. The taught metric. After repeat rate becomes the expansion’s headline metric, one store begins enrolling every walk-in customer in the loyalty program at checkout and logging a token second “purchase” through free-gift redemptions.
3. The quiet exit. A contribution-ranked analysis recommends closing the two stores in the markets where StyleCraft has historically spent the least on marketing. No one in the meeting represents the customers in those markets.

2.14.3 Extensions

Exercise 2.11 Rewrite the Request (Homework Submission or Optional Extension)

Rewrite each request below as an open analytical framing. For each, name the decision, the decision-maker you would confirm, and the deliverable you would propose.

1. “Build a dashboard proving that influencer marketing drives our growth.”
2. “Get me a number that shows the app is worth the investment.”
3. “Pull whatever data supports raising the free-shipping threshold.”
4. “Show the board that Gen Z loves the brand.”
5. “Just send me last month’s numbers.”

Exercise 2.12 Reflection Questions (Optional)

Finally, the following questions are intended to be thought-provoking rather than graded.

1. Think of an organization you know. Can you identify one report it produces that no decision depends on? What decision could it be attached to?
2. Have you ever been asked, in any setting, to find evidence for a conclusion that was already chosen? How did you, or how would you now, respond?

3. Which component of the analytic specification do you expect to be hardest to complete in real projects, and why?
4. When, in your judgment, should a marketing decision be automated, and what would make you comfortable delegating it?
5. What would make you trust a recommendation enough to act on it if you were the CMO?

2.15 Glossary of Terms

This glossary includes only the terms introduced in this chapter. Each definition is tied to the sources used in the chapter rather than added for decoration.

Analytic specification. A short written document, prepared before analysis begins, that records the decision to be supported, the decision-maker and deadline, the analytical questions, the unit of analysis, the data, the candidate methods, the deliverable, the verification plan, the limitations, and the planned use of AI assistance (adapted from Chapman et al., 2000; Provost & Fawcett, 2013).

Baseline. The comparison point against which a metric or result is judged, such as a prior period, a comparable group, a plan, a do-nothing alternative, or a stated expectation (adapted from Provost & Fawcett, 2013; Shmueli, 2010).

Constraint. Any condition, including time, data, budget, skills, legal or privacy limits, and organizational realities, that bounds what an analysis or decision can do (operational course concept).

Decision automation. The delegation of individual, repeated decisions to rules, algorithms, or models within boundaries and monitoring set by humans (Davenport et al., 2020; Huang & Rust, 2021).

Decision cadence. The frequency with which a decision recurs; a driver of the appropriate deliverable (operational course concept).

Decision quality. How well a choice was framed, informed, and reasoned at the time it was made, as distinct from the quality of the outcome that followed (Keeney, 1992; Kahneman et al., 2011; March, 1994).

Decision support. The use of data, analysis, models, and visualization to inform a specific managerial choice, while responsibility for the choice remains with a human decision-maker (Gorry & Scott Morton, 1971; Keen & Scott Morton, 1978; Davenport, 2006; Sharma et al., 2014).

Metric gaming. The reorganization of behavior around a measure once it becomes a target, degrading the measure's value as evidence (Campbell, 1979; Strathern, 1997).

Responsible framing. The practice of posing analytical questions as open questions, choosing metrics with awareness of gaming, and representing affected parties who are absent when

the decision is framed (adapted from Campbell, 1979; Strathern, 1997; Barocas & Selbst, 2016; Martin & Murphy, 2017).

Satisficing. Searching for an option, or an analysis, that is good enough given limits on time, information, and attention, rather than provably optimal (Simon, 1997).

Stakeholder. Any person or group that affects, or is affected by, the decision an analysis supports (Freeman, 1984).

Value of information. The worth of evidence, judged by its potential to improve the expected quality or value of the decision it informs, net of the cost of obtaining it (Howard, 1966; Keeney, 1992).

2.16 Further Readings

Students who want additional background may begin with the following readings. Decision-focused sources are listed first because this chapter approaches analytics from the decision side.

- Keeney (1992) for value-focused thinking and structuring decisions before analyzing them.
- Kahneman et al. (2011) for a practical checklist for reviewing evidence behind big decisions.
- March (1994) for an accessible treatment of how decisions actually happen in organizations.
- Keen and Scott Morton (1978) for the classic statement of decision support systems as aids to, not replacements for, managerial judgment.
- LaValle et al. (2011) for why analytics investments often fail to reach decisions.
- Davenport et al. (2020) and Huang and Rust (2021) for the augmentation and automation of marketing decisions.

2.17 References

- Anthony, R. N. (1965). *Planning and control systems: A framework for analysis*. Harvard University Graduate School of Business Administration.
- Barocas, S., & Selbst, A. D. (2016). Big data's disparate impact. *California Law Review*, 104(3), 671–732.
- Bezos, J. (2018). *2017 letter to shareholders*. Amazon.
<https://www.aboutamazon.com/news/company-news/2017-letter-to-shareholders>
- Campbell, D. T. (1979). Assessing the impact of planned social change. *Evaluation and Program Planning*, 2(1), 67–90. [https://doi.org/10.1016/0149-7189\(79\)90048-X](https://doi.org/10.1016/0149-7189(79)90048-X)
- Chapman, P., Clinton, J., Kerber, R., Khabaza, T., Reinartz, T., Shearer, C., & Wirth, R. (2000). *CRISP-DM 1.0: Step-by-step data mining guide*. CRISP-DM Consortium.

- Davenport, T. H. (2006). Competing on analytics. *Harvard Business Review*, 84(1), 98–107.
- Davenport, T. H., Guha, A., Grewal, D., & Bressgott, T. (2020). How artificial intelligence will change the future of marketing. *Journal of the Academy of Marketing Science*, 48, 24–42. <https://doi.org/10.1007/s11747-019-00696-0>
- Freeman, R. E. (1984). *Strategic management: A stakeholder approach*. Pitman.
- Google. (n.d.). *Google Ads automated bidding*. Google Ads Help. Retrieved July 20, 2026, from <https://support.google.com/google-ads/answer/10964872>
- Gorry, G. A., & Scott Morton, M. S. (1971). A framework for management information systems. *Sloan Management Review*, 13(1), 55–70.
- Howard, R. A. (1966). Information value theory. *IEEE Transactions on Systems Science and Cybernetics*, 2(1), 22–26. <https://doi.org/10.1109/TSSC.1966.300074>
- Huang, M.-H., & Rust, R. T. (2021). A strategic framework for artificial intelligence in marketing. *Journal of the Academy of Marketing Science*, 49, 30–50. <https://doi.org/10.1007/s11747-020-00749-9>
- Kahneman, D., Lovallo, D., & Sibony, O. (2011). The big idea: Before you make that big decision... *Harvard Business Review*, 89(6), 50–60.
- Keen, P. G. W., & Scott Morton, M. S. (1978). *Decision support systems: An organizational perspective*. Addison-Wesley.
- Keeney, R. L. (1992). *Value-focused thinking: A path to creative decisionmaking*. Harvard University Press.
- LaValle, S., Lesser, E., Shockley, R., Hopkins, M. S., & Kruschwitz, N. (2011). Big data, analytics and the path from insights to value. *MIT Sloan Management Review*, 52(2), 21–32.
- Lee, J.-Y., Gao, Z., & Brown, M. G. (2010). A study of the impact of package changes on orange juice demand. *Journal of Retailing and Consumer Services*, 17(6), 487–491. <https://doi.org/10.1016/j.jretconser.2010.08.003>
- March, J. G. (1994). *A primer on decision making: How decisions happen*. Free Press.
- Martin, K. D., & Murphy, P. E. (2017). The role of data privacy in marketing. *Journal of the Academy of Marketing Science*, 45, 135–155. <https://doi.org/10.1007/s11747-016-0495-4>
- Pfeffer, J., & Sutton, R. I. (2006). Evidence-based management. *Harvard Business Review*, 84(1), 62–74.
- Provost, F., & Fawcett, T. (2013). Data science and its relationship to big data and data-driven decision making. *Big Data*, 1(1), 51–59. <https://doi.org/10.1089/big.2013.1508>

- Sharma, R., Mithas, S., & Kankanhalli, A. (2014). Transforming decision-making processes: A research agenda for understanding the impact of business analytics on organisations. *European Journal of Information Systems*, 23(4), 433–441.
<https://doi.org/10.1057/ejis.2014.17>
- Shmueli, G. (2010). To explain or to predict? *Statistical Science*, 25(3), 289–310.
<https://doi.org/10.1214/10-STS330>
- Simon, H. A. (1997). *Administrative behavior: A study of decision-making processes in administrative organizations* (4th ed.). Free Press.
- Strathern, M. (1997). “Improving ratings”: Audit in the British university system. *European Review*, 5(3), 305–321. [https://doi.org/10.1002/\(SICI\)1234-981X\(199707\)5:3<305::AID-EURO184>3.0.CO;2-4](https://doi.org/10.1002/(SICI)1234-981X(199707)5:3<305::AID-EURO184>3.0.CO;2-4)