

Data Visualization for Communication and Decision-Making (Tableau II)

Audience, Hierarchy, Emphasis

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 builds the communicative half of data visualization, using StyleCraft's go/no-go expansion decision as its running material and Tableau as its instrument. It establishes the reader as the binding design constraint through audience analysis, a stated three-second reading test, and the choice among exploratory and explanatory forms. The KPI hierarchy then spatializes the metric discipline of Section 3.5 into primary, supporting, and diagnostic tiers, and layout converts that hierarchy into reading order. Narrative arc and the assertion title turn views that display quantities into views that state findings, at the cost of an obligation to verify every sentence against the marks beneath it; emphasis and decluttering make those sentences findable; interactivity is bounded to navigation, and on this artifact that bound turns out to exclude filter controls entirely. Uncertainty receives its own treatment, expressed in decision units and drawn as regions whose outlines are computed upstream. Two labs build the executive dashboard, test it on a reader, and pair it with a recommendation memo.

KEYWORDS

data visualization; audience analysis; dashboard design; KPI hierarchy; assertion titles; preattentive attributes; decluttering; uncertainty communication; usability testing; Tableau

VERSION AND DATE

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

SUGGESTED CITATION

Mendoza, J. (2026). Data visualization for communication and decision-making (Tableau II). In *Applied business analytics for marketing decision-making: Business analytics and data visualization* (Chapter 13, 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.

Tableau, Tableau Desktop, and Tableau Prep are trademarks of Salesforce, Inc. Google Colab is a product of Google LLC. ChatGPT is a product of OpenAI. Claude is a product of Anthropic. Gemini is a product of Google LLC. GitHub Copilot is a product of GitHub, Inc. NotebookLM is a product of Google LLC.

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, certified metric files, and figure sources for this chapter:

<https://github.com/jrmst102/businessanalytics>

GENERATIVE AI USE

Generative artificial intelligence and other AI-assisted tools were used in the research, writing, revision, and production of this chapter, including literature discovery, source organization, outlining, preliminary drafts, prose revision, support for code and analytical examples, and document formatting. These tools were used under the author's direction and are not credited as authors, researchers, or sources. The author determined the chapter's scope, learning objectives, methods, interpretations, and recommendations, and reviewed and approved all AI-assisted material: factual claims and citations were checked against the underlying sources rather than accepted from AI-generated summaries, and code and analytical outputs were tested or otherwise reviewed for accuracy. Responsibility for the accuracy, originality, and final form of this chapter rests entirely with the author. A fuller statement appears in the front matter of the complete guide.

Chapter Learning Objectives

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

1. Conduct an audience analysis for a specific visual deliverable, stating who will read it, what decision they must make, how long they will look, what they already know, and what would change their mind — and derive design consequences from each.
2. Apply the course's three-second reading test as a design constraint, predicting what a fresh reader will take away from an artifact before testing it, and revising the artifact rather than the reader when the prediction fails.
3. Distinguish exploratory from explanatory artifacts, choose among a dashboard, a report, and a data story for a stated decision and audience, and state the conditions under which a dashboard is the wrong deliverable.
4. Design an executive dashboard with an explicit KPI hierarchy — primary, supporting, and diagnostic — that spatializes the metric and KPI discipline of Section 3.5, so that a view's position and size encode its importance to the decision.
5. Lay out a dashboard for reading order, grouping, density, and whitespace, and justify each placement in terms of the sequence in which the reader must acquire the argument.
6. Construct a narrative arc across a set of views, write assertion titles that state what each view shows rather than what it plots, and verify every assertion against the underlying view before the artifact leaves the analyst's hands.
7. Use preattentive attributes deliberately to emphasize and de-emphasize, place annotation per Section 12.9 at the specific marks an assertion title depends on, apply data-ink discipline to remove what does not carry meaning, and explain why decluttering is a precondition for emphasis rather than an aesthetic preference.

8. Build interactivity that serves navigation and communication — filters with a stated scope, parameters used as navigation with titles that track them, tooltips, and dashboard actions — and identify the cases in which interactivity conceals the point rather than clarifying it.
9. Communicate the uncertainty established in Chapters 10 and 11 to an executive audience, constructing a prediction-interval ribbon from precomputed bounds at the forecast's own grain, and pairing the display with the asymmetric-cost reasoning of Section 2.7.
10. Pair a dashboard with a written recommendation that states the decision, the evidence, the uncertainty, and the conditions under which the recommendation would change.
11. Run a formative usability check on a visual artifact and a structured critique exchange with a peer, treating the reader's stated takeaway as data about the artifact, and document the resulting revisions.
12. Audit AI-generated titles, layouts, and recommendation prose against the views and evidence that must support them, and document the exchange per Appendix D.
13. State where emphasis becomes distortion in a persuasive artifact, defend a specific design choice against that standard, and identify the disclosures that keep a persuasive chart honest.

Chapter 12 built the analytical half of visualization and closed with an artifact designed to be argued with. The pre-read circulated on Monday, and by Wednesday morning it had done its job: three of the four readers replied, two of them at length, and the argument the December 2 meeting was going to have has now largely been had in writing. That is a better outcome than most analytical work achieves, and it is not the outcome this chapter needs. The pre-read is eight views, two of them faceted across seventeen store panels, every one carrying a footer recording source, window, filters, and grain — a document built for readers who have a week, a stake, and a reason to interrogate it. On December 2 the artifact will be read by a chief executive who has ten minutes, no prior exposure to the file, a real-estate committee waiting on her, and eight leases that price on Friday. Nothing about the pre-read survives contact with that reader, and the temptation the whole chapter is written against is to believe that it might, because the analysis inside it is correct and correct analysis feels as though it ought to travel. It does not travel. It has to be carried, and carrying it is a separate discipline with its own techniques, its own failure modes, and its own ethical boundary — the boundary between helping a decision-maker see what the evidence shows and helping her agree with what the analyst has concluded. Chapter 12 owned the mechanics by which charts mislead and declined to draw that boundary. This chapter draws it, in front of an audience, under a deadline, which is where it actually has to be drawn.

CONCEPT

What This Chapter Is Really About

Chapter 12 said that a chart is a machine for making a comparison easy. A communication artifact is a machine for making a *decision* easy, and the two sentences are not the same sentence, because a decision requires more than a legible comparison. It requires the reader to know which comparison matters most, to know what the comparison implies for the choice in front of her, to know how confident the analyst is, and to know what would have to be true for the recommendation to be wrong — and it requires her to acquire all four in the order that makes them usable, in the time she actually has. Every technique in this chapter is a way of supplying one of those four things without requiring the reader to do work she will not do. The KPI hierarchy tells her which comparison matters most, by putting it where the eye lands first.

The assertion title tells her what the comparison implies, by saying it in a sentence instead of leaving her to infer it from an axis label. Emphasis and decluttering make the sentence findable by removing everything that competes with it. The uncertainty view tells her how much to trust it. And the written recommendation tells her what to do and under what conditions to stop doing it. The uncomfortable part, and the reason this chapter needs an ethics section that Chapter 12 could defer, is that every one of those techniques also works when the analyst is wrong. A confident hierarchy, a fluent title, and a clean layout will carry a mistaken conclusion into a decision just as efficiently as a correct one — more efficiently, in fact, because a mistaken conclusion is usually the simpler story.

Chapter 12's verification asked whether the chart showed what the data contained. This chapter's verification asks a second question on top of that one: whether the reader took away what the chart actually supports. Only the reader can answer it, which is why this chapter ends with a test on a human being rather than a check on a number.

13.1 Marketing Decision Context: The Go/No-Go Meeting

The pre-read went out Monday at 4:40 p.m. By Wednesday morning three of the four readers had replied, and the replies are the reason this chapter exists.

The head of merchandising wrote first and wrote briefly: the composition view is the whole argument, she has been saying so for a year, and she wants it in front of the chief executive. The director of real estate wrote second and at length, and his note is the one that mattered. He had opened with the calendar-time facet — the one showing the eight Wave 4 stores as short low lines at the right edge — and had written most of a paragraph arguing that the newest stores were not performing before he reached the store-age facet on the same page and, in his words, "had to start the email over." He is now the expansion's most careful advocate, which is exactly what a good pre-read produces. The finance partner from Chapter 11 replied third, with a single question that the pre-read does not answer anywhere in its eight views: *at what point does one of these stores cover its own cost, and how many of the eight sites look like the ones that got there?*

That question is the December 2 meeting. StyleCraft's chief executive convenes the review at 9:00 a.m. and has, on the agenda circulated Friday, allotted it fifteen minutes, of which the

analyst should assume ten. Eight sites in the suburban and resort formats are under letter of intent — the fifth wave's flagship sites are not yet at that stage and are a separate decision on a separate benchmark; the letters of intent carry an eight-week option on any site StyleCraft does not execute on Friday, which is the fact that makes a conditional recommendation operationally available rather than rhetorical; the leases price at the end of the week; and the decision is a single binary with a conditional version available — commit to all eight, commit to none, or commit to a subset defined by a stated rule. The chief executive has not read the pre-read and will not. She has read the one-page summary the VP of Marketing wrote from it, she will arrive with the finance partner's question in her head because the finance partner reports to her, and she will make the call in the room.

The VP's brief to the analyst on Tuesday afternoon was four sentences long and is worth quoting because it is the actual design specification. One screen. Ten minutes. She should be able to say what she thinks after three seconds and then spend the rest of the time testing it. And she will ask what would have to be true for you to be wrong, so have that on the screen too, not in your head.

The first draft, built Tuesday night, does not meet any of those conditions, and the way it fails is instructive because it is the failure this chapter's students will reproduce on their first attempt. It was assembled the fast way, which is to say it was assembled out of the pre-read: the eight views of Lab 12.1 dragged onto a dashboard canvas, arranged four across and two down, each keeping the title Tableau generated from its field names — "SUM(Line Revenue) by Order Week and Metro," "Repeat-Order Share by Store Type" — each keeping the five-line analytical footer, each keeping its full legend, and the whole grid sharing a filter panel down the right side with seven controls on it. Everything in it is true. Every axis starts where it should, every grain is stated, every reference line is present and labeled, and the file it runs on is the certified extract that reconciles to the cent. It would survive any audit in Chapter 12, and shown to a reader for three seconds it communicates one thing, which is that a competent person has done a lot of work. Shown to the VP for rather longer than three seconds, it produced the question that names the problem: *what am I looking at first?*

There is no answer, because nothing on the canvas encodes an answer. Eight views of equal size in a regular grid assert that the eight matter equally, which is a claim the analyst does not believe and would not defend if asked; the largest visual object on the screen is the filter panel, which asserts that the reader's most important act is filtering; and the titles describe what each chart plots rather than what it shows, which leaves the reader to derive eight findings in ten minutes from material she has never seen. The draft is not wrong. It is unstructured, which for a ten-minute reader is a more expensive failure than being wrong would be, because a wrong chart gets caught in the room and an unstructured one just quietly fails to land and everyone blames the meeting for running long.

There is also a second problem, and it is the harder one. The finance partner's question does not have a clean answer, and answering it at all requires two numbers that the pre-read never needed and that no earlier chapter produced. The finance partner supplies the first as a planning

assumption: the annual fixed footprint of a new StyleCraft store in the suburban and resort formats — rent, staffing, and allocated overhead — is \$92,000, which is \$1,769 of contribution that a store must generate in an ordinary week to carry its own fixed cost. The certified extract supplies the second: the contribution margin on the suburban and resort store population, computed as $\text{line_revenue} \text{ less } \text{line_cost} \text{ over } \text{line_revenue}$, runs at 46 percent, so the weekly *revenue* at which such a store covers its footprint is approximately \$3,846. Against that declared threshold, the four mature suburban and resort stores first reached weekly coverage on a four-week average in weeks 43, 46, 51, and 57 of their operating lives — later than the weeks in which each first crossed the line once, which were 38, 42, 44, and 51, and the gap between the two is itself part of the answer, since a store whose two weeks are far apart crossed on a single week its neighbors did not repeat. The eight Wave 4 stores, all of them suburban or resort, are seventeen weeks old. Two of them are tracking above the range their predecessors occupied at the same store age, five are inside it, and one, Scottsdale, has been below it since week six. Section 13.12 states where every one of those figures comes from, which population it describes, and requires the student to derive it before any title uses it, because a number that arrives in a chapter without a derivation is a number a student cannot check.

Chapter 10's forecast supplies the last of the material, and it has to be rebuilt before it can be used. Chapter 10 produced a six-month planning forecast of chain *revenue*, because that was what the merchandising plan consumed. A lease decision spans a planning year and turns on contribution rather than revenue, so this chapter reruns Chapter 10's model-selection procedure on the daily contribution series — the same validation design, not the same winner assumed to transfer — and reports the fiscal year the leases will live in. Two properties of the result matter in the room. It is an **outlook rather than a forecast**: July through November have already happened by December 2 and are carried as actuals, so only the seven months from December onward have a range around them. And it is stated in **contribution**, which means the interval needs no margin applied to it before it can be compared with anything. StyleCraft's FY2027 chain contribution is therefore \$0.51 million already banked plus a central estimate of \$0.74 million still to come, or \$1.25 million, with the part still open running from \$0.62 million to \$0.86 million. The width of that range is not a technicality: end to end it is worth roughly two mature stores' annual contribution. The honest answer to the chief executive's question is therefore conditional, and conditional answers are the ones that communicate worst — they are the ones that get compressed in the retelling, that get heard as *the analyst is not sure*, and that tempt the analyst into stating something cleaner than the evidence supports in order to be useful. Section 13.9 is about that temptation and Section 13.15 is about where giving in to it stops being a communication choice and becomes an ethical one.

So the commission is a single artifact with an unusual specification: one dashboard, readable in three seconds at the top level and interrogable for ten minutes below it, carrying a recommendation that is conditional without being evasive, accompanied by a written memo that says what the screen cannot. Building it is Lab 13.1. Testing whether it works — on an actual reader, before the meeting rather than during it — is Lab 13.2, and it is the chapter's verification

theme, because the only evidence that a communication artifact communicates is what a person who reads it says afterward.

13.1.1 Opening Case Questions

Keep these questions in mind while reading, and return to them after completing the labs.

- The first draft contains every finding in the pre-read and is correct in every particular. Write one sentence stating what it fails to do, and one sentence explaining why "the information is all there" is not a defense for a ten-minute reader.
- The chief executive will look at the screen for roughly three seconds before forming a first impression. Name the one sentence you would want her to form, and then name the design decision — position, size, color, or title — that you would use to produce it.
- The finance partner's question has a conditional answer: some of the eight sites resemble stores that reached weekly coverage and some do not. List two ways to communicate a conditional recommendation without either overstating confidence or sounding as though you have no recommendation, and predict which section of this chapter each comes from.
- A dashboard shows eight gray lines of equal weight on a single set of axes. You want the reader to find one particular line without searching for it. Name the single visual property you would change, state why changing two properties at once would work less well, and predict which section of this chapter explains the answer.
- The pre-read was built to be disagreed with. The dashboard is built to support a decision. Name one thing the dashboard must add that the pre-read deliberately withheld, and one thing the pre-read carried that the dashboard should drop — and state, for the second one, what is lost by dropping it.

13.2 Know Your Reader

Section 13.1 described a dashboard that fails for a reason no audit of its charts would detect. This section supplies the diagnosis, and the diagnosis is that the draft was designed against the data rather than against a person. Chapter 12's standard for a view was internal — does the encoding support the comparison, is the axis honest, is the grain what the metric requires — and every one of those questions can be answered by the analyst alone at her own screen. This chapter's standard is external. It depends on a reader with a finite amount of attention, a specific decision, and prior knowledge the analyst does not control, and it cannot be answered by inspection. It can only be predicted and then tested, which is why the reader has to be specified before anything is built.

DEFINITION

Audience Analysis

Audience analysis is the explicit statement, made before an artifact is designed, of who will read it, what decision or action it must support, how much time and attention the reader will give it, what the reader already knows and believes, what technical vocabulary can be assumed, and what evidence would change the reader's mind. Its output is a set of design constraints rather than a description: each element of the analysis rules some design choices in and others out, and an artifact that cannot be traced back to the analysis is being designed against the analyst's own knowledge rather than the reader's.

Source: Adapted from Knaflitz (2015) and Few (2006).

In other words, the reader is a specification in the sense Section 2.5 established, and audience analysis is what makes the specification writable. The five questions that do the most work are worth stating with StyleCraft's answers attached, because the answers are what dictate the design. Who is reading? A chief executive who is the sole decision-maker, with four colleagues in the room who have already read the pre-read and will not need convincing of the analysis, only of the recommendation. What decision? Whether to commit to eight leases, in whole, in part, or not at all, by Friday. How long? Ten minutes total, of which the first three seconds determine what she thinks she is looking at and the remaining minutes are spent testing that impression. What does she know? That StyleCraft opened seventeen stores in four waves; that the suburban customers spend more per order; and, because the VP's one-page summary said so, that the newest stores are young rather than weak. And what would change her mind? A comparable that failed, or a footprint figure that moved. What does she not know? Every number in the pre-read, the grain of any of them, and the vocabulary of Part II — she has never heard the phrase *prediction interval* used in a meeting she was chairing, and an artifact that requires her to learn it in the first minute has spent a tenth of its budget on vocabulary.

That last exchange is the one students find hardest to accept, and it is worth being precise about what it does and does not license. It does not license removing the uncertainty; Section 13.9 will insist on the opposite. It licenses expressing the uncertainty in a form that does not require a definition, which is a design problem rather than a compromise. It also does not license writing down to the reader. A chief executive who has run a retail business for a decade knows a great deal that the analyst does not, and the failure mode of oversimplification — a screen so reduced that it insults the reader's judgment and gives her nothing to test — is as real as the failure mode of density. What the analysis licenses is spending the reader's attention on the things only the analyst can supply.

The operational form of all of this is a test that takes three seconds to run and is the single most useful instrument in this chapter. It is a course heuristic rather than a standardized procedure, and it is worth saying so plainly, because the technique students will meet in the usability literature is normally a five-second first-impression test (Doncaster, 2014). Three seconds is a deliberately harder constraint, chosen because an executive dashboard on a projector

in a room with an agenda gets less attention than a web page does, and because a test an artifact barely passes teaches less than one it must be designed to pass.

DEFINITION

The Three-Second Test (Course Heuristic)

The three-second test is the practice of showing a finished visual artifact to a representative reader for approximately three seconds, removing it, and asking what the artifact was about. Inspired by first-impression testing and by the glanceability principles of interface design, it operationalizes the fact that a reader forms a first impression within a few fixations and then interprets everything subsequent in light of it; the test therefore measures not whether the artifact can be understood but what it communicates before the reader has decided to work at it. This guide uses a three-second exposure as an intentionally demanding executive-dashboard constraint; the formal first-impression procedure in the usability literature is normally run at five seconds (Doncaster, 2014), and either duration is defensible provided it is stated.

The artifact passes when the reader's stated takeaway matches the analyst's intended takeaway, and the correct response to a failure is to revise the artifact, since the reader's response is the measurement rather than the error.

Source: Course heuristic developed for this guide, informed by Doncaster (2014), Few (2006), and Krug (2014).

In other words, three seconds is not a claim about how long executives read. It is an instrument for isolating what the artifact's visual hierarchy is doing, stripped of the reader's willingness to be helpful. Given thirty seconds, a cooperative colleague will find the finding in almost any dashboard, and the analyst will learn nothing about her design; given three, the colleague reports what the design actually pushed at her, which is the measurement of interest. Krug's (2014) argument about web pages transfers and is worth carrying: readers do not study an interface, they glance at it, form a plausible story, and act on the story, and a designer's job is to make the plausible story the true one.

Two rules of use keep the test honest. The first is that the analyst must predict the answer before running it, in writing, which is this chapter's version of the discipline the guide has enforced since Section 1.7 — a prediction converts the test from a reaction into evidence, and the gap between predicted and actual takeaway is the finding. The second is that the reader must be representative rather than convenient. A classmate who built the same dashboard last week is not a test of anything, because she knows where to look; the useful reader is someone who knows the business context and has not seen the artifact, which in a workplace means a colleague from another team and in this course means a peer working on different data.

Finally, a caution about what the test cannot do, since students routinely overextend it. The three-second test measures the top of the hierarchy and nothing else. It says whether the primary message lands; it says nothing about whether the supporting evidence is legible, whether the interactivity is discoverable, or whether the reader can find the answer to a question she thinks of in minute six. Those require the longer protocol of Section 13.10, which is a different instrument

aimed at a different failure. An artifact that passes the three-second test and fails the ten-minute one is a poster rather than a dashboard, and the December 2 meeting needs both properties in the same object.

13.3 Exploratory vs. Explanatory Artifacts

Section 13.2 specified the reader. This section takes the next decision, which comes before layout and before any chart is drawn, and which the first draft skipped entirely: what kind of object is being built. Chapter 12 distinguished visualization's two jobs and assigned this chapter the second one, but "communication artifact" is not a form. It is a category containing at least three forms with different structures, different reading protocols, and different failure modes, and choosing among them is the highest-leverage decision in the chapter, because a well-built object of the wrong kind cannot be repaired by refinement.

DEFINITION

Exploratory and Explanatory Artifacts

An exploratory artifact is designed to let a reader investigate a body of data by her own path, supporting questions the builder did not anticipate; its structure is a space to be navigated, and its success criterion is coverage of the questions a reader might legitimately ask. An explanatory artifact is designed to deliver a specific finding along a path the builder chose, supporting a question the builder has already answered; its structure is a sequence, and its success criterion is that the intended understanding is acquired in the time available.

The distinction is not the same as the analysis-communication distinction of Section 12.2: an artifact may be built for an external audience and still be exploratory, and the most common design error in business visualization is building an exploratory object for a reader who needed an explanatory one.

Source: Adapted from Knaflitz (2015) and Segel and Heer (2010).

In other words, the question is whether the reader is looking for something or being shown something, and it is answered by the decision rather than by the data. Three forms follow, and a fourth object travels with all of them; Table 13.1 places them against the readers and decisions they serve.

Table 13.1

Choosing the Artifact from the Decision and the Reader

Artifact	What it is	Reader and time	Use when	Fails when
Dashboard	A single screen of coordinated views, usually refreshing, with a fixed set of questions designed in	A recurring reader with minutes, returning on a schedule	A decision or a monitoring task recurs on a known cadence and the same questions are asked each time	The decision is one-time; the reader is new to the material; the questions cannot be anticipated
Report	A linear document combining prose, charts, and tables, read at the reader's pace	A reader with an hour, reading alone, who may need to re-read	The evidence is long, the reasoning must be inspectable, and the reader will need the document again later	The reader will not read it; the decision happens in a meeting nobody prepared for
Data story	An ordered sequence of views with narrative connective tissue, delivered live or as a designed scroll	A captive audience with a fixed span, in a room or on a screen	A specific finding must be transferred once, in order, to people who have not seen the evidence	The reader needs to explore; the finding is genuinely contingent on the reader's own question
Recommendation memo	Prose stating the decision, the evidence, the uncertainty, and the conditions for reversal	A decision-maker, before or after the meeting	A decision requires a record of what was recommended and on what basis	Substituted for the visual evidence rather than paired with it

The pre-read of Chapter 12 was, in these terms, an exploratory artifact for an external audience — deliberately, because its readers' job was to find its weaknesses. The December 2 artifact is explanatory and, on the surface, a dashboard. The interesting part is that a strict reading of Table 13.1 says a dashboard is the wrong form here, because the decision is one-time and the reader is new to the material, and the analyst should notice this rather than shrug at it. What resolves the conflict is that the artifact is not being handed to the chief executive to read alone; it is being presented, live, in a room, by the person who built it, with a memo attached. That makes it a data story delivered on a dashboard's canvas, with the dashboard's interactivity reserved for the questions the room asks after the story lands — and it is the reason Lab 13.1 builds a single screen with a strict reading order rather than the free-navigation grid the first draft produced. Naming this explicitly is not pedantry. It settles half a dozen later decisions: it is why the filter

panel disappears entirely, why the views are unequal in size, why the titles carry sentences, and why there is a memo at all.

DEFINITION

Data Story

A data story is an ordered sequence of views, each carrying a stated point, connected so that the sequence itself constitutes an argument — typically moving from the situation the reader already accepts, through the complication the evidence introduces, to the resolution the analyst recommends. Its defining property is that the order is chosen by the author and is load-bearing: rearranging the views changes what the audience concludes, which is what distinguishes a data story from a collection of correct charts about the same subject.

Source: Adapted from Segel and Heer (2010) and Knaflitz (2015).

One further note closes the section, and it is the one the Business Analytics in Practice section will find running through every organization that has deployed self-service tools at scale. The dashboard is over-selected. It is what stakeholders ask for by name, it is what tool vendors sell, and it is what analysts default to because it looks like the most work for the least commitment — a dashboard implies no particular conclusion, which is comfortable. The discipline this chapter asks for is to answer the question in Table 13.1's *Use when* column honestly before opening a canvas. A one-time decision made in a meeting by a reader who has never seen the data is a data story with a memo attached, and building it as a dashboard because a dashboard was requested is how an analyst spends two days producing something nobody opens twice.

13.4 KPI Hierarchy on a Dashboard

Section 13.3 settled what kind of object is being built. This section supplies its internal structure, and the structure is the direct spatial expression of a discipline students have owned since Chapter 3. Section 3.5 distinguished metrics from KPIs, and the distinction was that a KPI is a metric elevated by a decision: of the dozens of quantities an analysis can produce, a small number are the ones a specific choice actually turns on. That elevation has been carried in prose ever since. On a dashboard it has to be carried in space, because space is what a reader decodes first and fastest, and a canvas on which every view is the same size has flattened the hierarchy and thereby thrown away the analyst's most important judgment.

DEFINITION

KPI Hierarchy

A KPI hierarchy is the ordering of a deliverable's metrics into tiers by their proximity to the decision — a primary tier of one or two measures on which the decision turns, a supporting tier that explains or qualifies the primary, and a diagnostic tier that answers the follow-up questions the primary provokes — and the expression of that ordering in the artifact's visual prominence, so that size, position, and emphasis correspond to decision relevance. Because readers allocate attention by prominence before they read any label, a hierarchy that is stated in the analyst's head but not in the layout is not communicated at all.

Source: Adapted from Few (2006) and the metric and KPI roles of Section 3.5.

In other words, the first design act on a canvas is to decide what the reader should look at first, second, and third, and the second design act is to make the canvas say so. StyleCraft's case makes the tiers concrete, and building them requires going back to the decision rather than to the pre-read. The decision is whether eight leases are signed. The quantity that decision turns on is not revenue, is not average order value, and is not the repeat-order share that dominated the pre-read's argument; it is how long a new store takes to reach, on a four-week average rather than in a single week, the point where its contribution covers its own fixed footprint. Everything else on the screen is there to explain that number, to qualify it, or to answer the objection it invites. Table 13.2 sets the tiers out in the order the reader will meet them, and it is written to be filled in before a canvas is opened rather than read once.

Table 13.2

The KPI Hierarchy of the Go/No-Go Dashboard

Tier	What it must establish	The StyleCraft measure	Visual treatment	Reader's question it answers
Primary	Whether a store clears its own fixed cost, and when	First four-week-average coverage: the first store-age week in which the four-week average of weekly contribution reaches the \$1,769 that carries the \$92,000 annual fixed footprint, for each mature suburban and resort store	Largest object on the canvas, upper left, with the threshold drawn, the four-week-average coverage week labeled on each of the four lines, and the earlier one-week crossing marked more faintly	"Do these stores cover their own cost, and when?"

Tier	What it must establish	The StyleCraft measure	Visual treatment	Reader's question it answers
Primary	Where the eight new stores sit against that path	Wave 4 store ramps overlaid on the range their predecessors occupied at the same store age	Second-largest, adjacent to the primary, sharing its vertical scale	"And are the new ones on track?"
Supporting	Why the suburban economics differ at all	Average order value by metro, computed at order grain, with the chain average drawn	Half-height, beneath the primaries, first in reading order	"What is actually different about these customers?"
Supporting	Whether that difference is basket composition or price	Category revenue share by metro, with occasionwear identifiable	Half-height, second in the supporting band	"Is it a mix story or a pricing story?"
Supporting	What the planning year is worth, and how much of it is still open	The FY2027 chain contribution outlook — actual through November, forecast thereafter — with its 80 and 95 percent regions on the forecast months and a labeled horizon-total block, annotated with the open range in store-equivalents	Half-height, rightmost in the supporting band	"How much of this is a guess?"
Diagnostic	The known weakness in the case	Repeat-order share by store type against the chain share, with the suburban gap labeled	Small, lower band, de-emphasized but present	"What are you not telling me?"
Diagnostic	The store that does not fit the pattern	The one Wave 4 store tracking below the predecessor range, named, with its ramp	Small, lower band, adjacent to the above	"Which sites would you not sign?"

Three properties of that table deserve comment, because each one is a decision students tend to skip. The first is that the primary tier contains two views rather than one, and that this is close to the limit. A hierarchy with four primaries has no primary, and the discipline of forcing the count down to one or two is where the analyst discovers what she actually believes the decision turns on. The second is that the diagnostic tier contains the case's two weakest facts — the suburban repeat-order share of 24 percent against a chain share of 47 percent, and the Scottsdale store that

has tracked below the predecessor range since week six — and that they are on the screen rather than in a backup slide. Putting them in the diagnostic tier is not a concession; it is what makes the primary tier believable, and Section 13.15 will argue that it is also the ethical requirement. An executive who discovers a weakness herself concludes that the analyst either missed it or hid it, and both conclusions are expensive. The third is that every tier is defined by a question the reader will ask, in the order she will ask it, which is what turns a hierarchy into a sequence and hands the next section its material.

One boundary belongs here. The tiers govern prominence, not inclusion, and the temptation the tier structure creates is to solve a crowded canvas by demoting rather than by cutting. A view that has no tier does not become a small view; it comes off the screen and goes into the appendix of the memo, where it remains available to anyone who asks and stops competing for the ten minutes. Section 13.7 gives the general form of that rule, and it is the hardest habit in this chapter to acquire, because everything on the first draft was work the analyst did.

13.5 Dashboard Layout

The hierarchy of Section 13.4 is a claim about importance. This section makes the claim visible, and it does so through four instruments — reading order, grouping, density, and whitespace — that together convert a list of tiers into a canvas a reader traverses in the intended sequence. The section is deliberately practical, because layout is where the previous three sections either become a design or remain an intention.

Begin with reading order, which is the instrument that carries the sequence. Readers of left-to-right languages enter a rectangular canvas at the upper left and move rightward and downward, and the eye's path across a screen is shaped by that habit and by the size and contrast of what it encounters — a large, high-contrast object anywhere on the canvas will capture the entry fixation regardless of position, which means position and prominence must agree or they will fight. The practical rule that follows is simple and is violated constantly: place the primary tier at the upper left and make it the largest object, place the supporting tier along the path the eye takes next, and place the diagnostic tier last in that path. The go/no-go dashboard therefore runs the coverage view across roughly the upper-left 40 percent of the width, the Wave 4 ramp comparison immediately to its right at about 35 percent and sharing its vertical scale, the three supporting views as a half-height band beneath them, and the two diagnostics as a smaller band at the bottom. That arrangement is not a grid; it is a sentence, and its reading order is the argument of Section 13.6 rendered as geometry.

Grouping is the second instrument, and it works by proximity and enclosure. Objects placed close together are read as related and objects separated by whitespace or a rule are read as distinct, and this happens pre-consciously and cannot be overridden by a label. The two consequences that matter on a dashboard are that views belonging to the same tier should be visually adjacent and share their formatting, and that views sharing a scale or an axis should be aligned so that the shared scale is visible as such. The Wave 4 ramp view sits directly beside the cost-coverage view for exactly this reason: the two carry the same measure on the same vertical

range, and placing them adjacent with aligned axes lets the reader compare across identical non-aligned scales — one rung down the perceptual hierarchy of Section 12.4 rather than several, which is the best available given that they cannot share a single frame.

Density is the third and is where executive dashboards most often fail in the direction opposite to the one students expect. The failure is not too few views; it is too many marks per view and too many views per screen, each individually defensible. Miller's (1956) observation that people hold only a handful of items in working memory is invoked constantly in design writing and is often overinterpreted, but the practical form of it survives the caveats: a reader who must hold eight findings simultaneously in order to reach a conclusion will not reach it, and a canvas of eight equal views demands exactly that. The counter-discipline is to state, before laying anything out, the number of things the reader must hold at once — for the go/no-go dashboard it is two, the cost-coverage path and the new stores' position on it — and to build the layout so that everything else is available in sequence rather than simultaneously.

Whitespace is the fourth, and it is the one students cut first and should cut last. Empty space is not wasted space; it is the mechanism by which grouping and hierarchy are expressed, and a canvas with generous margins around its primary view is telling the reader where to look as surely as an arrow would. The most common mistake in Tableau specifically is to leave the sizing decision to the software without deciding what consumption condition the artifact is built for. Tableau offers Fixed size, Range, Automatic, and device-specific layouts, and each is correct for a different situation (Tableau, 2026e). Automatic stretches every object to consume whatever window it is opened in, which is the right behavior for a dashboard published to the web and opened on unknown displays and the wrong behavior for an artifact whose hierarchy is expressed in relative area and which will be projected at a known resolution in a known room. The go/no-go dashboard is built Fixed at the projection resolution because the room is known; a monitoring dashboard read on phones by eight field managers should use device-specific layouts instead, and the choice between them is a statement about where the artifact will be read rather than a preference.

Figure 13.1

The Go/No-Go Dashboard, as a Wireframe

Figure 13.1 The Decision Before the Dashboard

StyleCraft December 2 Decision: Execute 8 Suburban and Resort Leases (in Full, in Part, or Not at All) Recommendation due by Friday, December 4		
 DECISION CONTEXT Decision-maker: StyleCraft Chief Executive Choice: Execute eight signed letters of intent for new stores (in full, in part, or not at all) Deadline: Friday, December 4 Consequence: Eight lease commitments will incur fixed costs regardless of opening timing.	 DECISION CRITERIA Coverage: Stores should cover their fixed footprint within a reasonable time. Trajectory: New stores should track with or above the precedents early. Economics: The portfolio must generate adequate contribution after fixed costs. Risk: Downside must be understood and bounded. Capacity: Operations must be able to support the planned openings.	 KEY FACTS RECORD (from Finance Partner) Fixed footprint per Suburban store \$92,000 per year = \$1,769 per week Contribution margin (Suburban & Resort) 46% <i>(Line Revenue – Line Cost) / Line Revenue</i> Weekly revenue required to cover footprint \$3,846 = \$1,769 / 46%
 AUDIENCE ANALYSIS (5 LINES) Who reads it: StyleCraft Chief Executive. Decision: Execute eight suburban and resort leases in full, in part, or not at all. How long: ~3 seconds for top message; ~90 seconds for details. What she knows: StyleCraft model, markets, prior performance, fixed footprint, Wave 4 plan. What would change her mind: Evidence that stores will not cover in reasonable time, downside outweighs upside, or capacity is insufficient.	 INTENDED TAKEAWAY (ONE SENTENCE) <i>Proceed with five of the eight suburban and resort leases now, hold three, and re-evaluate in eight weeks based on the Wave 4 trajectory.</i>	 BOUNDARIES AND ASSUMPTIONS Data: Certified StyleCraft orders and lines through November 30, 2025. Focus: Suburban and Resort formats only. Time windows: Mature comparison through Week 64; Wave 4 observed through Week 17; forecast Weeks 18–78. Excludes: Flagship format and international stores. Assumes: No material changes in pricing, cost structure, or operating model. Uncertainty: Forecast interval reflects model and planning uncertainty.
 GUIDING QUESTION	<i>Given what we know now, will these stores cover their fixed footprints in a reasonable time and earn adequate contribution after fixed costs, with risk we are willing to accept?</i>	

CONCEPT

Position Is a Claim About Importance

Every element on a canvas occupies a location, and the reader decodes location before she decodes anything else — before the title, before the axis, before the numbers. This means an analyst who has not decided where things go has still made a claim about what matters, and has made it by accident, usually in the order the views happened to be built. The habit this section asks for is to treat placement as an assertion that must be justified in the same way a title must: for every view on the canvas, be able to state why it sits where it sits, in terms of the tier it belongs to and the point in the reader's sequence at which she needs it. Two consequences follow immediately, and both are uncomfortable.

The first is that the upper-left position is scarce and cannot be given to the view the analyst is proudest of; it belongs to the view the decision turns on, which is frequently a plainer object. The second is that a view for which no justification can be stated does not belong on the canvas at all — and the honest test of whether the hierarchy is real is whether anything was removed to make room for it.

Two structural mechanics finish the section, since Tableau supplies them and the labs will use them. Layout containers — horizontal and vertical — are the mechanism by which the tiers are actually held in place, and building a dashboard out of nested containers rather than floating objects is what makes the layout survive a change in screen size or a late addition. And a dashboard needs a title band and a source band, which are the communicative descendants of Section 12.9's footer discipline: the title band carries the decision the screen supports and the date, and the source band carries the extract, the window, and the filter state in small type. What changes from Chapter 12 is not whether the provenance appears but how much of it appears at each level. Chapter 12 put five lines under every view because every view was going to be interrogated individually; a dashboard carries the provenance once, for the screen, with per-view grain statements moved into tooltips where the reader who wants them can find them and the reader who does not is not taxed. That is a legitimate compression and not a concealment, and the line between the two is exactly what Section 13.15 has to define.

13.6 Storytelling with Data

Section 13.5 arranged the views in the order the reader will traverse them. This section supplies the reason the order is that order, and the sentences that make each stop on the path mean something. The word *storytelling* has been abused enough in analytics marketing to make graduate students suspicious of it, and the suspicion is healthy, because a great deal of what travels under the name is decoration applied to a deck. What the term denotes here is narrow and technical: an argument with a stated order, in which each view carries one claim, and in which the sequence is doing work that no individual view does.

The structure that has proven most durable is three-part and predates data visualization by a considerable margin. It opens with the situation, which is the state of affairs the audience already accepts and which exists to establish common ground rather than to inform; it introduces the complication, which is the fact that makes the situation inadequate and which is where the analysis enters; and it closes with the resolution, which is what the analyst recommends doing about it. StyleCraft's version writes itself once the KPI hierarchy is fixed. The situation is that StyleCraft opened seventeen stores in four waves and that the suburban customers spend more per order than the urban ones — which everyone in the room believes and which the composition and order-value views establish in ten seconds. The complication is that revenue per order is not the quantity the lease decision turns on; the quantity is the time it takes a store to reach, on a four-week average rather than in a single week, the weekly contribution that covers its own fixed footprint, and on that measure the mature stores' four-week moving averages did not first reach

that threshold until somewhere between weeks 43 and 57, which is longer than anyone in the room has been assuming. The resolution is that the eight sites are not one decision but three: the ones whose matched comparable reached four-week-average coverage, the one whose matched comparable did not, and the ones with no qualifying comparable at all.

Notice what the sequence accomplishes that no single view does. It converts the finance partner's question from an objection into the organizing question of the artifact, which is generally the right move when a skeptical reader has already voiced one — a story built around the strongest objection is far more persuasive than one built around the strongest evidence, because the audience is not left holding an unaddressed doubt while the analyst talks. Segel and Heer (2010) surveyed narrative visualization across journalism and analysis and identified this same tension as the genre's organizing axis: every such artifact sits somewhere between author-driven ordering, which is efficient and directive, and reader-driven exploration, which is thorough and slow. The go/no-go dashboard sits well toward the author-driven end for its first minute and moves toward the reader-driven end afterward, which is a design position rather than a compromise, and it is implemented through the interactivity of Section 13.8.

The sentence-level instrument is the title, and it is where most of the chapter's persuasive power and most of its ethical risk are concentrated.

DEFINITION

Assertion Title

An assertion title is a chart title that states the finding the view supports, written as a complete sentence with a subject and a verb — "The four mature suburban and resort stores' four-week moving averages first reached the weekly fixed-cost threshold in weeks 43, 46, 51, and 57" — in contrast to a label title, which names the fields displayed — "Weeks to Weekly Coverage by Store." Because a reader acquires the title before decoding the marks and interprets the marks in light of it, an assertion title determines what the view is understood to show; the technique therefore carries an obligation that a label title does not, which is that the sentence must be verifiable against the view directly beneath it, in the view's own units, without additional evidence.

Source: Adapted from Knaflitz (2015) and Kosslyn (2006).

In other words, the assertion title is the insight statement that Chapter 5 distinguished from a mere observation, written above the view that proves it — which is why it is the single highest-return technique in this chapter and the single most dangerous, and both facts have the same cause: the reader believes it. Three rules keep it honest, and they are what Lab 13.2 checks. The first is that the sentence must be true of the view and not merely true — a title reading "Suburban stores are our best growth opportunity" may be defensible somewhere in the analysis and is not defensible above a chart of weeks-to-coverage, and a reader who verifies the sentence against the marks will find that it cannot be done. The second is that the verb must be one the evidence licenses, which is Section 11.10's evidence hierarchy arriving in the title bar: *rose*, *is*, *reached*, and *differs* are descriptive verbs a view can support; *drives*, *causes*, *proves*, and *will* are claims about mechanism or the future that a descriptive view cannot carry, however strong the

underlying analysis. The third is that a number in the title must be a number in the view. Titles that assert a quantity the reader cannot find in the marks — "Suburban expansion adds \$1.2 million" over a chart showing weeks — are the most common way a fluent title outruns its evidence, and they are especially common when the title is drafted by an assistant that has been given the chapter's conclusion but not the view's data.

Table 13.3 works the transformation for the go/no-go dashboard's principal views, and the exercise of writing all three columns is itself the discipline: the label title is what the analyst has; the assertion title is what she believes; and the *What must be visible* column is the check that stops the second from drifting past the first.

Table 13.3

Label Titles and Assertion Titles for the Dashboard's Principal Views

View	Label title (what it plots)	Assertion title (what it shows)	What must be visible in the view for the assertion to be earned
Coverage path	Contribution 4wk Avg by Store Age and Store	"The four mature suburban and resort stores' four-week moving averages first reached the weekly fixed-cost threshold in weeks 43, 46, 51, and 57"	All four stores drawn as four-week averages, the \$1,769 weekly threshold drawn, the four-week-average coverage week labeled on each line, the earlier one-week crossing marked, and the zero-order weeks present rather than skipped
Wave 4 ramps	Contribution 4wk Avg by Store Age and Store	"Seven of the eight new stores are within or above the range their predecessors occupied at the same age"	The predecessor range as a region whose edges move week by week; exactly eight lines in the frame and no others, all of them Wave 4 and all suburban or resort; a shared vertical scale; store age on the horizontal axis
Order value	Average Order Value by Home Metro	"Suburban customers order \$156 at a time against a chain average of \$108"	Both figures as marks, the chain average drawn as a reference line, order grain stated
Mix	Category Revenue Share by Metro	"Occasionwear is roughly three times the share of suburban revenue that it is of urban"	Both metros' shares, occasionwear identifiable, shares summing to 100 percent
FY outlook	Monthly Chain Contribution with Prediction Intervals	"The FY2027 outlook is \$1.25 million of chain contribution, and the range still open on the remaining seven months spans about 1.8 mature stores' annual contribution end to end"	The monthly line, the seam line where reporting ends with no shading to its left, and both interval regions to its right; a labeled horizon-total block splitting actual from forecast, since neither total is readable from the monthly marks; and the store-equivalent annotation

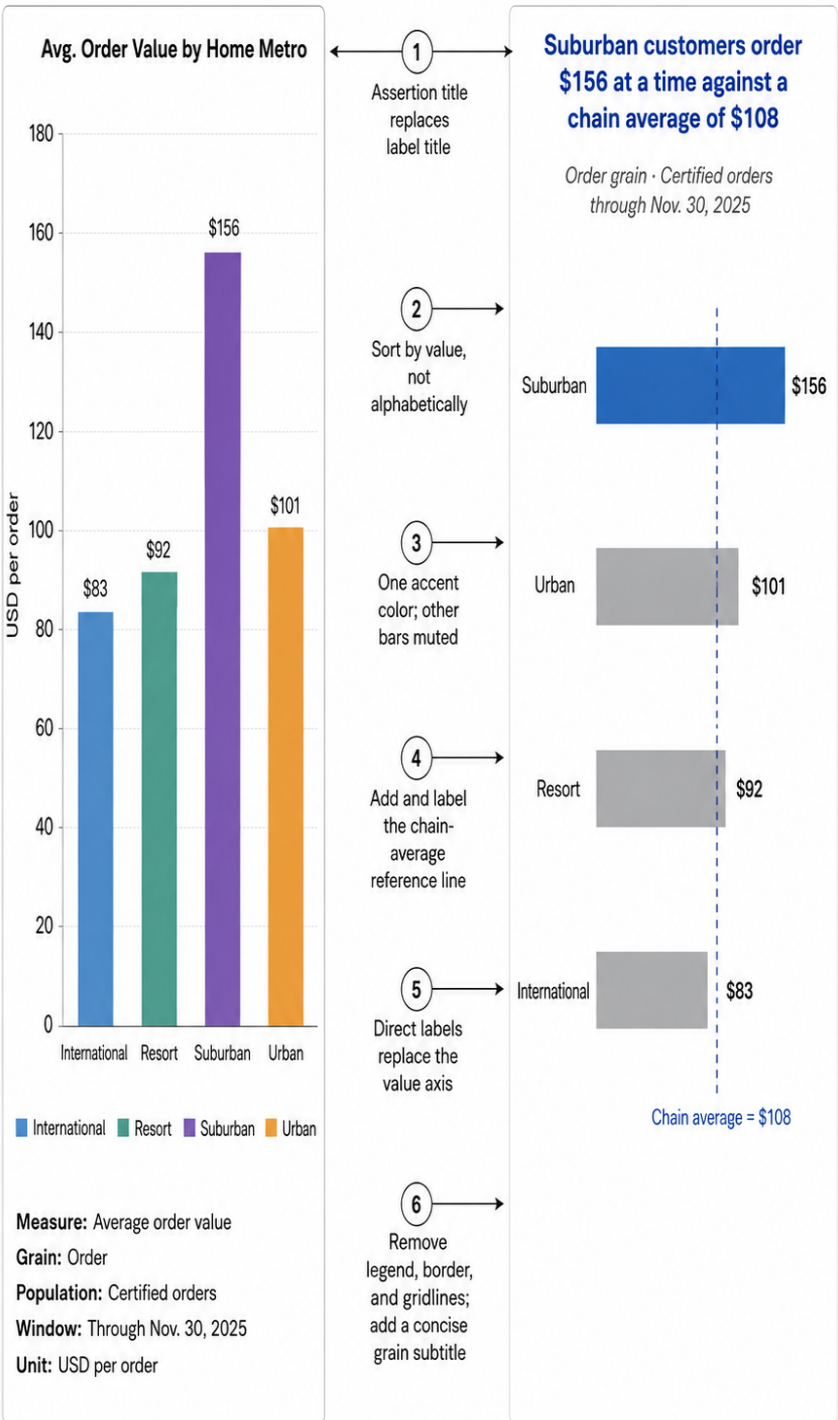
View	Label title (what it plots)	Assertion title (what it shows)	What must be visible in the view for the assertion to be earned
Repeat share	Repeat-Order Share by Store Type	"Suburban repeat-order share is 24 percent against a chain average of 47 percent"	Both figures, the chain reference line, store types sorted

Two general points close the section. The first is that the assertion title changes what a chart is for, and therefore what must be on it. The moment a title says "seven of the eight," the view acquires an obligation to show eight things distinguishably, which the pre-read's version did not, and the title has consequently forced a design change rather than merely labeling one — this is the productive direction of the technique and the reason titles should be drafted early rather than added at the end. It also disciplines the wording: the same view supports "within or above the predecessor range" and does not support "at or above the band," which reads as though all seven sat at the upper boundary or beyond it when in fact two are above the range and five are inside it. The second is that the title is where the analyst's voice is most exposed and should therefore be the last thing checked before the artifact ships. Section 13.11 makes that check a routine, and Lab 13.2 makes a peer run it, because the author of a sentence is the worst available judge of whether it says more than she can prove.

Figure 13.2

The Same Data as an Analytical View and as an Explanatory View

Figure 13.2 The Same Data as an Analytical View and as an Explanatory View



13.7 Emphasis, De-Emphasis, and Decluttering

Section 13.6 gave each view a sentence. This section makes the sentence findable, and it does so with two operations that are usually taught separately and are in fact one operation performed in two directions. Emphasis makes an element stand out. Decluttering removes the elements it would otherwise have to stand out against. Neither works alone: emphasis applied to a busy canvas produces a busy canvas with one loud thing in it, and decluttering without emphasis produces a clean canvas with no entry point.

The mechanism underneath emphasis is perceptual and has been measured, which lets this section do for emphasis what Section 12.4 did for accuracy.

DEFINITION

Preattentive Attributes

Preattentive attributes are visual properties that the visual system processes in parallel across much of the visual field, before focused attention is directed anywhere, so that an object differing from its neighbors in one such property can be detected rapidly and with little increase in search time as the number of neighbors grows. The principal attributes available in business charts are color hue, color intensity, size, position, orientation, shape, enclosure, and added marks. The effect is strongest under favorable conditions — a single distinctive feature, homogeneous distractors, adequate contrast — and it degrades with target-distractor similarity, with heterogeneous displays, when the target is defined by a conjunction of two attributes, and when several attributes vary at once.

It is therefore a design resource with conditions attached rather than a guarantee of constant-time detection.

Source: Adapted from Treisman and Gelade (1980) and Healey and Enns (2012).

In other words, a reader finds a red bar among gray bars without searching, finds a large mark among small ones without searching, and does search — slowly, and with errors — for the bar that is both red and large in a field where some bars are red and some are large. Treisman and Gelade's (1980) feature-integration experiments established the pattern and the conjunction penalty; Healey and Enns (2012) reviewed four decades of the subsequent work in visualization terms, including the ways the effect weakens as displays become less obliging. The practical consequence for a dashboard is a budget rather than a technique: preattentive emphasis works because it is scarce, and a canvas that emphasizes six things has emphasized nothing, because the reader must now search the emphasized set. Table 13.4 lists the attributes with the use each is suited to and the failure each produces when overspent.

Table 13.4

The Emphasis Toolkit

Attribute	What it does well	Use on the go/no-go dashboard	Failure when overspent
Color hue	Marks one category as the subject; separates a small number of groups	The four mature suburban and resort stores in the accent hue; every other store gray	Every series in a different hue, so the reader must consult a legend to know what she is looking at
Color intensity	Orders a quantity; recedes a background element	The predecessor ribbon as a pale fill behind the Wave 4 lines	Used for an unordered field, asserting an order (see Section 12.8)
Size	Signals importance at the layout level; carries a rough magnitude	The primary tier's views drawn largest on the canvas	Size differences within a view read as magnitude claims that were not intended
Position	Carries importance in reading order; carries value precisely	Primary tier upper left; supporting tier beneath	Prominence and position disagree, so the eye is pulled away from the intended entry point
Enclosure	Groups elements as belonging together; isolates a region of a chart	A light box around the eight Wave 4 lines in the ramp view	Boxes around everything, which restores the undifferentiated grid
Added marks	Directs attention to specific places, and loses its force as their number grows	One labeled annotation at each store's four-week-average coverage week on the coverage view — four in all, and nothing else annotated	Annotations on every mark, which become a second layer of clutter
Direct labels	Removes the legend lookup entirely	Store names at the end of the four mature lines	Labels on all seventeen lines, which is a thicket

The second operation is subtraction, and its governing idea is Tufte's.

DEFINITION**Data-Ink and Decluttering**

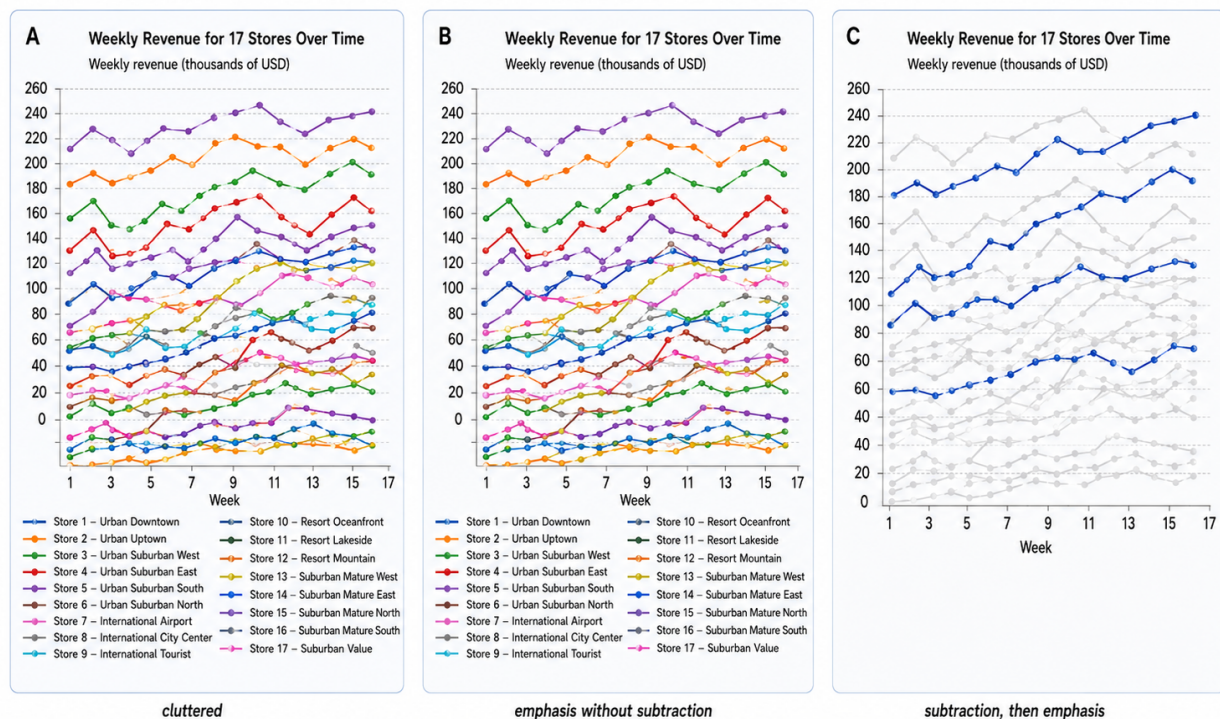
Data-ink is the portion of a chart's ink that represents data values and would be lost if the data changed; non-data-ink is everything else — gridlines, borders, backgrounds, tick marks, legends, redundant labels, and decoration. Decluttering is the systematic removal or de-emphasis of non-data-ink and of duplicated data-ink, undertaken not for austerity but to raise the contrast between what carries meaning and what does not, so that emphasis has something to work against. The discipline is a ratio rather than a minimum: elements that help a reader interpret correctly, including reference lines, direct labels, and annotation, are retained even though they encode no data value.

In other words, the rule is not "remove everything" but "remove everything that is not doing a job you can name," and the qualification matters because the naive version of data-ink maximization produces unreadable charts. Section 12.9's reference lines are non-data-ink and are mandatory; the footer is non-data-ink and is an ethical instrument; a gridline that lets a reader read a value off a long axis is non-data-ink doing real work. What comes off the go/no-go dashboard is the material the first draft inherited from the analytical version and never questioned: the heavy borders around every worksheet, the gray background panels, the tick marks at every increment, the axis titles that repeat what the assertion title already says, the four legends replaced by direct labels, the three views whose findings the memo carries better than a chart does, and all seven filter controls. That last removal is the largest single improvement to the draft, and it is a removal of function rather than of decoration, which is why Section 13.8 has to justify it.

Figure 13.3

Decluttering Is What Makes Emphasis Possible

Figure 13.3 Decluttering Is What Makes Emphasis Possible



CONCEPT

Decluttering Is Deleting Your Own Work

The reason decluttering is hard is not that analysts lack taste. It is that every element on a first draft is there because somebody put it there, usually for a reason that was good at the time, and removing it feels like discarding effort — which it is. The two habits that make it survivable are procedural rather than aesthetic. The first is to declutter against the artifact's stated sentence rather than against a general principle: with the assertion title of Section 13.6 written down, every element on the canvas can be asked whether it helps a reader believe that specific sentence, and elements that do not have a defensible answer come off without further argument. The second is to move rather than to delete, at least at first.

A view that does not earn its place on the screen goes into the memo's appendix; a filter that does not earn a control goes into the workbook the analyst brings to the room; a footnote that does not earn space in the source band goes into the tooltip. Nothing is lost, the reader's ten minutes are recovered, and the analyst is spared the psychological cost that makes this step the one most often skipped. What is not permitted is the third option, which is to shrink everything until it fits — a canvas of twelve small views is not a decluttered canvas of twelve views, it is an illegible one.

13.8 Interactivity as Navigation and Communication

Sections 13.6 and 13.7 built an artifact that says one thing clearly. This section takes up the property that makes a dashboard different from a printed page, and it argues for a use of that property considerably narrower than the tool invites. Interactivity in this chapter is a means of navigation and of communication: it lets a reader move through an argument at her own pace, ask a follow-up the analyst anticipated, and see a detail without cluttering the canvas. It is not, in this chapter, a means of computation. Parameters that recompute a metric, what-if controls, threshold sliders that change what a view is measuring, and set actions that build cohorts on the fly are analytic instruments rather than communicative ones; they belong to the advanced visual-analytics work that follows this guide, and a student who reaches for them here should notice that she is rebuilding the exploratory artifact Section 13.3 already ruled out.

The organizing idea is Shneiderman's (1996) mantra, which has aged well: overview first, zoom and filter, then details on demand. Read as a design sequence rather than as a slogan, it says that the reader should be able to see the whole before she is asked to choose a part, that filtering should narrow a view she already understands rather than construct one she does not, and that detail should arrive when requested rather than occupy the canvas permanently. Heer and Shneiderman (2012) later cataloged the interaction types that implement it, and the four that matter for a marketing dashboard are worth naming with their communicative jobs attached. A filter narrows the population under discussion and answers "what about just the suburban stores?" A highlight action marks the same entity across several views at once and answers "where else does Scottsdale appear?" A tooltip supplies the row-level or grain-level detail the canvas cannot carry and answers "what exactly is this mark?" And a navigation action moves the

reader from one screen to another and answers "show me the evidence behind that." Each of those is a question the analyst predicted; that is the test of whether the interaction belongs, and it is what the definition below turns into a design obligation.

DEFINITION

Dashboard Action

A dashboard action is an interaction configured on a dashboard by which activity in one view — a hover, a selection, or a menu choice — changes another view, by filtering it, by highlighting a matching entity within it, or by navigating the reader to a different sheet or an external resource. A highlight action operates on matching field values, so it can only mark an entity in a target view that actually contains that entity; a highlight configured from a source view onto a target whose data does not include the selected value does nothing, silently.

An action's communicative purpose is to answer a follow-up question the builder anticipated without giving that question permanent space on the canvas; its design test is therefore that the analyst can name the reader's question the action answers, that the target views contain the values the action will pass them, and that a reader who never triggers it still acquires every finding the artifact must communicate.

Source: Adapted from Heer and Shneiderman (2012), Shneiderman (1996), and Tableau (2026b).

That second condition is where most configured actions quietly fail, and the go/no-go dashboard supplies two examples of it, one of which is a direction rather than a target. A highlight on `store_name` cannot "trace one store across the argument" into the cost-coverage view at all, because that view contains only the four mature suburban and resort stores and shares no store name with Wave 4; it is left out of the action rather than included and inert. The subtler point is which of the two remaining views is the source. Sourcing from the eight-store ramp view onto the single-store outlier view fires for one selection in eight and does nothing for the other seven — not an error message, just no response, which a reader interprets as a broken dashboard rather than as an empty match. Sourcing from the outlier view onto the ramp view fires for every selection it can produce, because the one store it contains is one of the eight the target holds. The general rule is that an action should run from the smaller value set to the larger one, and Lab 13.1 Part C builds it that way. Table 13.5 states each interaction's communicative case, its concealing case, and the disclosure it requires, and its *Conceals when* column is the one to read carefully, because every entry in it describes something an analyst does by accident rather than by design.

Table 13.5

Interactivity That Communicates and Interactivity That Conceals

Interaction	Communicates when	Conceals when	The disclosure it requires
Filter control	It narrows a population the reader already understands, the default state is the honest one, and the set of worksheets it applies to is stated	The default is a restricted state the reader assumes is complete; the control is the largest object on the screen; it silently changes a view whose title asserts a chain-wide figure	The current filter state and the worksheets affected, visible as text on the canvas, not only in the control
Highlight action	It links the same entity across views that all contain that entity, so a reader can trace one store through the part of the argument where it appears	It is the only way to see a comparison the reader needs, so a passive reader never sees it; or its target views do not contain the field values, so it appears to do nothing	Nothing; but the comparison must also be legible without the interaction, and the action must run from the smaller value set to the larger so that every selection it can produce has a match
Tooltip	It carries grain, definition, and provenance for readers who want them	It carries the qualification that makes the headline defensible	Anything load-bearing moves out of the tooltip and onto the canvas
Navigation action	It moves a reader from a summary to its supporting evidence on request	It buries a finding two clicks deep so the summary reads cleaner than the evidence supports	A visible cue that the deeper view exists
Parameter used as navigation	It swaps which of several equivalent views is displayed — one metro at a time, say	It changes the evidence on display in a way the title does not track	The title must be built dynamically so that it remains true in every state the control can produce

The last row is the one that produces real accidents, and it is worth stating as a rule because it is the intersection of this section with Section 13.6. An assertion title is a sentence about a view; a parameter changes the view; therefore an assertion title on a parameterized view must change with it, or the analyst has shipped a screen that is true in one state and false in the others. Tableau supports inserting a field or parameter value into a worksheet title, so the sentence can be assembled in a calculated field and inserted (Tableau, 2026f); Lab 13.1 Part C does it, and the failure is common enough in circulated workbooks that it is worth looking for in anyone else's.

Filters carry a quieter version of the same problem, and it is worth being explicit because Tableau's scoping options make it easy to get wrong. A filter applied to all worksheets using the same data source will change every view built on that source, including views whose titles assert chain-wide quantities and views whose reference lines were computed on a different population (Tableau, 2026c). Scoping is the first tool to reach for: a store-type filter belongs on the views

that are about physical store performance and not on the order-value, mix, outlook, or repeat-share views, whose assertion titles quote chain-wide figures that must not move when the reader narrows the store set. Reference values entered as constants from a certified metrics file stay fixed under filtering by construction; a reference computed from the view changes with the marks, so filtering it moves the benchmark as well as the bars and the title becomes false. But scoping is not always sufficient, and the next paragraph is the case where it is not.

Sometimes no scope is safe, and then the control comes off. That is what happens to both of the go/no-go dashboard's surviving filters, and the two cases fail differently, which is why both are worth following. A date-window control fails on the *titles*: the order-value, mix, and repeat-share views quote full-window figures in their sentences and draw full-window constants as reference lines, so narrowing the window moves the marks and leaves four numbers and two reference lines standing behind them, unchanged and now wrong. A `store_type` control fails on the *rows*. Scoped to the two ramp views and the outlier view it can empty the coverage view entirely when the reader deselects a format; it can change the number of lines in a frame whose title says "seven of the eight"; and — because the predecessor region's outline rows carry no store type at all — it can delete the region while leaving the subtitle that describes it. Scoping narrows which worksheets a filter touches; it does not keep a sentence true, and it does not protect rows for which the filtered field is null (Tableau, 2026c).

Lab 13.1 Part C therefore ships an executive artifact with **no filter controls on it**, and states the reasoning rather than leaving it to be inferred. The two alternatives were to make six titles dynamic — spending the reader's ten minutes on questions she has not asked — or to remove the controls and declare the state in words instead. A screen with no controls makes no claim about its own completeness, so the source band has to carry the population, the window, the exclusions, and the sentence *no filters applied*. All seven controls live on an analytical tab in the same workbook, one click away for the analyst and invisible to the room. The general rule this produces is worth carrying: a control that can falsify a sentence on the same screen is not a convenience, and "the reader can always set it back" is not an answer, because the reader does not know what it was set to.

Two cautions close the section. The first is that interactivity is invisible until it is discovered, and readers in meetings do not explore. An executive who has ten minutes will click nothing; a colleague who opens the workbook on Thursday will click a little; the analyst who built it will click everything and will therefore badly overestimate how much of her design the audience will ever see. The working rule is that every finding the artifact must communicate has to be visible in the default state, and interactivity may only add depth to a finding that has already landed. The second is the one the first draft violated: the presence of many controls is itself a message, and the message is "this is a tool for you to use," which is the wrong message for a ten-minute reader making a one-time decision. Cutting the go/no-go dashboard's filter panel from seven controls to none is a communication decision rather than a simplification, and all seven survive on an analytical tab in the workbook the analyst brings to the room, where they answer questions if questions come.

13.9 Communicating Uncertainty to Executives

The artifact now has a hierarchy, a sequence, a set of sentences, and a controlled amount of interaction. This section addresses the property those instruments most easily destroy. Chapter 10 produced a forecast with a prediction interval and insisted that the interval was the honest part; Chapter 11 produced a measured lift with a confidence interval and a break-even line beneath it; Chapter 9 produced a threshold derived from asymmetric costs. All three of those results have a range attached, and all three are about to be presented to a reader who wants a number. What happens in that collision is the subject here, and the guide's position is stated at the outset so that the argument can be read against it: the interval is what makes the recommendation defensible, not what weakens it, and an analyst who strips it has not made her advice more useful but has transferred the risk from her own credibility to her employer's balance sheet.

Start with why the stripping happens, because it is not usually cowardice. Hullman (2020) surveyed ninety visualization authors and interviewed thirteen influential designers about why uncertainty so rarely appears in published work, and documented a consistent tension: authors describe uncertainty as valuable and omit it anyway, because it is hard to represent well, because they fear an audience will read a range as incompetence, and because a distribution takes space and explanation that a point estimate does not. Those are real constraints and they are the design problem, not an excuse to skip it. Spiegelhalter et al. (2011) surveyed the available forms for communicating uncertainty about the future to non-technical audiences and reached a conclusion this section adopts: the choice of form matters more than the choice of statistic, and forms that let a reader see the range as a set of possible worlds outperform forms that ask her to interpret an abstraction.

Three practical moves follow, and the go/no-go dashboard uses all three.

The first is to express the interval in decision units rather than in statistical ones, and to say what the interval is before converting it. StyleCraft's FY2027 outlook is \$1.25 million of chain contribution — \$0.51 million already banked between July and November, plus a central estimate of \$0.74 million for the seven months still to come — with the part still open running from \$0.62 million to \$0.86 million. That sentence is accurate and it means nothing to a chief executive in the three seconds it will get. Two conversions make it mean something, and the chapter's position is that a dashboard should carry both, because each answers a different question. The end-to-end width of the range still open is \$240,000 of contribution, and \$240,000 divided by the \$137,000 of annual contribution a mature store produces is 1.8 stores. Relative to the central estimate, the same range is approximately plus or minus \$120,000, which is roughly plus or minus 0.9 mature stores. The end-to-end figure describes how far apart the two ends of the plausible range are; the plus-or-minus figure describes how far the outcome might sit from the number the plan is being built on, and quoting only the first overstates how uncertain the central estimate is. The annotation the labs build says both: *the range still open spans \$240,000 end to end — about 1.8 mature stores' annual contribution, or roughly plus or minus 0.9 around the estimate.*

Notice what is absent from that arithmetic, because its absence was designed in. There is no margin in it. The forecast is stated in contribution and the denominator is stated in contribution, so the conversion is one division rather than a multiplication followed by a division. That is not a convenience; it is what keeps an unestablished claim out of the chain. Converting a chain-level *revenue* interval into store-equivalents would require multiplying it by a margin, and the only margin available is the 46 percent computed on the suburban and resort population — which is not the chain's margin, and assuming it is would bury an unsupported assumption inside an arithmetic step where no reader would ever find it. Selecting and fitting the model on contribution removes the step and the assumption together, which is the general principle: **when a conversion requires an assumption you cannot establish, change what you forecast rather than disclosing the assumption and proceeding.** Disclosure is the second-best answer and is available when the first is not. This is the *error in decision units* point of Table 8.5 applied to a range rather than to a point, and it is also where a careless analyst introduces a new error, which is almost always a units error, so the conversion is shown on the canvas in steps rather than asserted as a result, and the denominator is named as well as used.

The second is to draw the range as a ribbon rather than as a bar, and to build the ribbon from bounds that vary across the horizontal axis. A shaded interval region around a forecast line reads as a set of plausible futures and is decoded quickly; error bars on a bar chart are decoded with systematic biases that Correll and Gleicher (2014) measured, including a within-the-bar bias in which points inside a bar are judged more likely than points outside it at the same distance from the mean, and a tendency to read the interval's endpoints as a hard boundary rather than as a level of a continuous distribution. Their contribution was not only the diagnosis: they showed that gradient and violin encodings improved judgments on the same tasks, which is the reason this chapter prefers a graded region to a discrete marker. Nested regions — an 80 percent ribbon inside a 95 percent ribbon, each labeled — do not display the full probability density, and it is worth being exact about what they do and do not do: they communicate graded levels of uncertainty, showing that outcomes nearer the center are more concentrated than outcomes near the edges, without claiming to show how much more. The labs build the two-interval version because the contrast between the inner and outer region is what teaches the distinction, and because building it correctly forces the student to bring both sets of bounds into Tableau at the forecast's own grain rather than asking the software for a band.

That last point is the one that separates a correct uncertainty view from a plausible-looking one, and it deserves its own statement, because it is where the first draft of this dashboard failed. A reference band in Tableau is defined between two values at the table, pane, or cell scope, which produces a horizontal region across a pane rather than a region whose upper and lower edges move with the horizontal axis (Tableau, 2026a). A prediction interval that widens with the horizon is not a horizontal region. It has a different lower and upper value in every month, which means the bounds must exist as fields at month grain before Tableau is asked to draw anything, and the region must be constructed from those fields rather than requested from the Analytics pane. Tableau's own forecasting feature will draw a widening band, but it draws the band

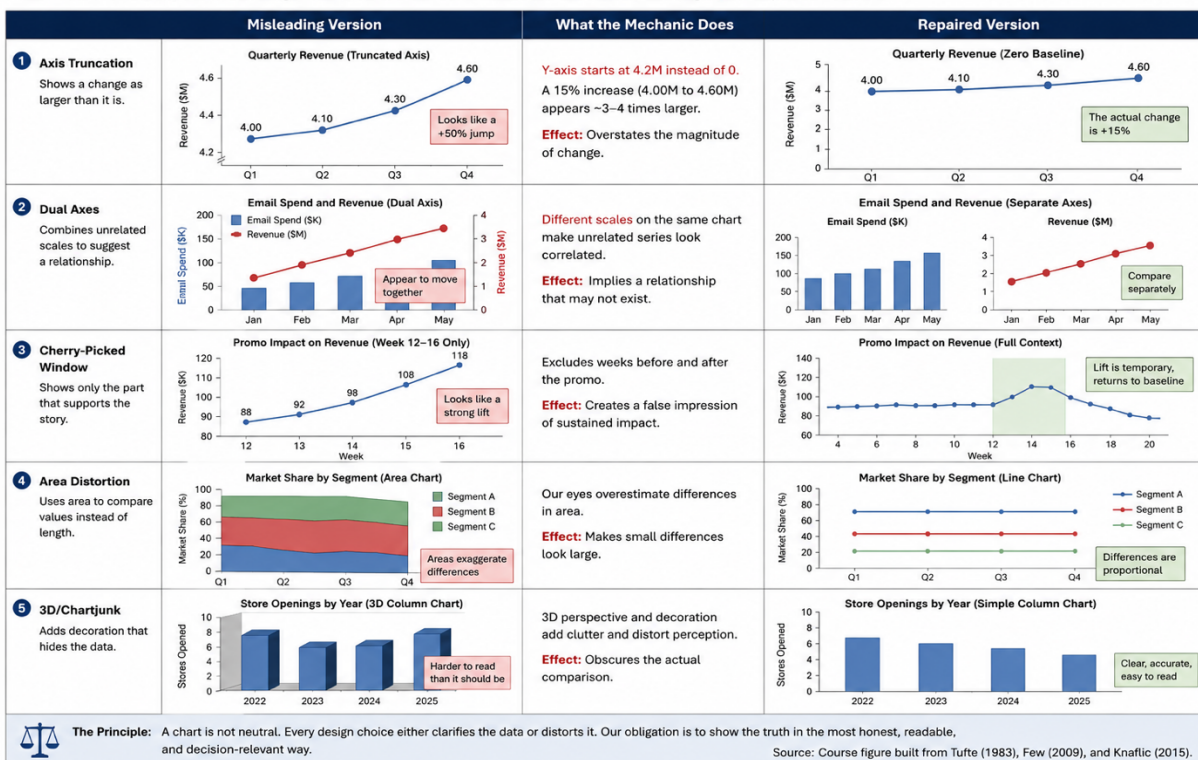
belonging to the forecast Tableau itself computed, at the single prediction interval currently configured (Tableau, 2026g); this chapter deliberately imports an externally produced forecast and displays two intervals at once, so neither convenience is available. The construction the labs use is a **polygon**: the region's outline is emitted upstream as an ordered list of vertices — along the lower bound left to right, back along the upper bound — and Tableau is asked only to connect them. Two nested intervals become three non-overlapping regions rather than two layered ones, so nothing can hide behind anything. Lab 13.1 Part D builds it, and the same construction draws the predecessor range in Part B, so the technique is learned once and applied twice. What it will not let a student do is the thing the tool makes tempting: drop two bound measures side by side on Rows and expect a filled region between them. Separate continuous measures on Rows occupy separate axes, stacking measures inside one pane requires Measure Values and Measure Names (Tableau, 2026h), and every segment stacked that way shares one Marks card and therefore one opacity, because opacity is a property of the card rather than of a measure (Tableau, 2026i).

Figure 13.4

How an Interval Ribbon Is Actually Built

Figure 12.4 The Five Misleading-Chart Mechanics (with StyleCraft Examples and Repairs)

Each row shows (left) the misleading version, what the mechanic does to the reader, and (right) the repaired version.



The third practical move is to pair the range with the decision rule it implies, which is where Chapter 2's asymmetric-cost discipline finally earns its keep in a room. A range alone invites paralysis; a range with a stated threshold and an asymmetry invites a decision. The go/no-go version is the finance partner's own arithmetic returned to him, and getting it right requires care, because the two sides are easy to state in units that do not compare. The downside of *declining* a site that would have worked is not the \$137,000 of annual contribution a mature store produces; that figure is contribution *before* the fixed footprint the site would also have carried, so the quantity actually forgone is the net — roughly \$137,000 less \$92,000, or about \$45,000 a year, before build-out, taxes, and working capital. The downside of *signing* a site that never reaches coverage is likewise not the full \$92,000; it is the part of the footprint the store's own contribution fails to cover, plus whatever is irreversible — the lease term, the fit-out, and the cost of exiting — and that last component is the one with no upper bound anywhere on this dashboard. Stated in comparable units the asymmetry runs toward caution rather than balance: the annual regret from a bad signing is bounded below by the unrecovered footprint and can exceed it by the exit cost, while the annual regret from a good site declined is closer to \$45,000 and is partly recoverable, since a declined site can often be revisited and a signed one cannot be unsigned. That is the reasoning that makes *hold* the risk-appropriate default for a site whose comparable evidence is absent or unfavorable, and it is why the recommendation is conditional on the ramp evidence rather than a blanket judgment about the market. It is also why the paragraph on the screen states the comparison qualitatively and shows both subtractions: neither side is known precisely, and a screen that set \$92,000 against \$137,000 as though they were comparable would be doing arithmetic the case does not support. Stating that reasoning is what turns the uncertainty view from a hedge into the argument's load-bearing element.

CONCEPT

Hedging Without Paralysis

There is a form of honesty that is indistinguishable, in a meeting, from having nothing to say. It sounds like this: "The results are suggestive but the interval is wide and there are limitations, so it is hard to draw a firm conclusion." Everything in that sentence may be true and it is nonetheless a failure of the analyst's job, because the decision is being made on Friday whether or not the analyst helps, and an executive who receives no recommendation does not therefore make no decision — she makes it on the basis of whoever in the room is most confident, which is a worse process than the one the analysis was supposed to improve. The alternative is not false confidence.

It is a recommendation stated with its conditions attached, which has a reliable four-part shape: *here is what I recommend; here is the evidence it rests on; here is the range I cannot narrow, in units you care about; and here is what would have to be observed for me to change the recommendation.* The fourth part is the one that converts hedging into usefulness, because it turns the uncertainty into a monitoring plan rather than a disclaimer.

"We recommend signing the five sites matched to comparables that reached four-week-average coverage, holding the site matched to Scottsdale and the two with no qualifying comparable, and revisiting all three in eight weeks when the Wave 4 stores reach the age at which the mature stores' paths had already separated" is a conditional recommendation that a chief executive can act on Friday, and it is more honest than a confident number, not less.

13.10 From Dashboard to Decision

The artifact is built and its uncertainty is stated. Two things remain before December 2, and both are about closing the gap between what the analyst intends and what actually happens in a room. The first is the memo. The second is the test.

Take the memo first, because students consistently treat it as an afterthought and it is the deliverable most likely to survive the meeting. Charts and marks do not by themselves constitute a durable recommendation. A dashboard can certainly contain one — a callout, an annotation, a decision rule written on the canvas, a text object stating what the analyst advises — and the go/no-go dashboard's title band does exactly that. What a dashboard cannot do is carry the conditions, the scope, and the reversal criteria that make a recommendation actionable and auditable, because those are prose and there is no arrangement of marks that says *unless*. The pairing this guide asks for is therefore a one-page written recommendation accompanying every decision-supporting visual artifact, structured as the four-part shape of Section 13.9's concept box and carrying, in an appendix, the views that came off the canvas during decluttering. Three properties make it work. It is written before the meeting, which forces the analyst to discover whether she can actually state the recommendation — a recommendation that will not survive being written down will not survive being questioned. It is left behind, which means it is what circulates to the people who were not in the room and what is consulted in March when someone asks what was decided and why. And it states the reversal conditions, which is what makes the analyst accountable in a way a screen never is, since a screen that was wrong can always be reread as merely displaying.

Now the test, which is this chapter's verification theme and the reason the labs end where they do.

DEFINITION

Usability Test

A usability test of a visual artifact is a structured observation in which a representative reader, who has not seen the artifact, is given a task or a question and observed while attempting it, with the analyst recording what the reader looks at, what she says, what she asks, and what she concludes — and refraining from explaining, guiding, or defending. Its purpose is to measure the artifact's communication rather than the reader's competence: a point at which the reader hesitates, misreads, or asks a question the artifact should have answered is evidence of a possible mismatch among the artifact, the task, and the audience, and is investigated as a candidate defect in the artifact rather than dismissed as a property of the reader.

Because a first impression cannot be had twice, a reader who has seen an artifact can no longer supply a fresh first-impression measurement of it, though she remains available for later task-completion and revision testing.

Source: Adapted from Krug (2014) and Nielsen and Landauer (1993).

In other words, this is predict-then-verify with a human being as the instrument, and it is the only form of verification in this chapter that the analyst cannot perform alone. The protocol is short. Predict, in writing, what the reader will say the artifact is about, which question she will ask first, and where she will hesitate. Show the artifact for three seconds and record the takeaway. Show it again for ninety seconds with one task — "tell me what you would decide and why" — and record what she reads, in what order, and what she asks. Then stop, and grade the predictions. What makes the protocol hard is not the mechanics but the discipline of silence: the reflex to explain the chart the moment the reader misreads it destroys the measurement and is nearly irresistible, which is why the lab makes the analyst write her predictions down first, so that there is something at stake in shutting up.

How many readers are needed is a question with a well-known answer that is frequently overstated in both directions. Nielsen and Landauer (1993) modeled the discovery of usability problems as a function of the number of test participants and found sharply diminishing returns, which is the origin of the widely repeated "five users" heuristic. The model's parameters were estimated on particular systems and user groups, and it does not establish five as a universal number for every interface, audience, or task, nor does it license testing once and declaring victory. What this course asks for is smaller and more defensible: a formative usability check on one or two peers before a meeting, which is enormously better than none, with the second reader seeing the revised artifact rather than the same one. A check of that size finds the gross hierarchy and language failures and does not validate that the artifact will work for executives — it is a way of catching what the analyst cannot see, not evidence that the artifact is finished. The analyst who tests nobody is relying on the assumption that she is a representative reader of her own work, which she is definitionally not.

The classroom version of the test is a structured exchange rather than an observation, and it is the instrument Project #2 runs.

DEFINITION

Critique Protocol

A critique protocol is a structured procedure for obtaining and giving feedback on a visual artifact, in which the reviewer proceeds in a fixed order — first reporting what the artifact appears to say, then asking what decision it is meant to support, then identifying specific elements that impede the reading, and only then proposing changes — and in which the author listens without defending until the sequence is complete. The order is the mechanism: reporting the takeaway before learning the intent yields an uncontaminated measurement, and separating observation from prescription keeps the exchange focused on the artifact's effect rather than on the reviewer's taste.

The protocol's value is that it makes the feedback usable. Unstructured critique of visual work reliably degenerates into preference — the reviewer would have used a different color, would have put the chart on the left, does not like the font — and preference is unfalsifiable and therefore unhelpful. A protocol that requires the reviewer to state the takeaway first produces a fact the author cannot argue with, and a protocol that requires her to name the element responsible before proposing a fix produces a change the author can evaluate. It also protects the reviewer, since "I read this as saying the new stores are failing" is an observation about herself rather than an accusation about the author's work, and observations of that kind are much easier to deliver to a peer.

One closing note connects the section back to the decision it serves. The purpose of all of this — the memo, the test, the protocol — is not to make the analyst's artifact better in the abstract. It is to reduce the probability that a correct analysis produces a wrong decision, which is the specific failure mode Part III exists to prevent and which is invisible in every measure of analytical quality that Chapters 4 through 11 supplied. A model with a good holdout error that leads to a bad lease decision has failed. That failure happens in the last two feet, it happens fast, and the only defense against it is to find out, before the meeting, what a person who is not the analyst actually sees.

13.11 AI as a Communication Assistant

Chapter 12 observed that a chart's mechanics and its meaning are produced by the same act, which is why an AI-drafted chart cannot be audited the way an AI-drafted model can. Communication work sharpens that observation into something more uncomfortable. An assistant asked to write a chart title, a dashboard headline, or a paragraph of recommendation prose produces language, and language is the medium in which a claim is made — which means that in this chapter the assistant is not drafting an instrument whose output the analyst will interpret. It is drafting the interpretation itself. The gap between what a view supports and what a sentence about that view asserts is the entire subject of Sections 13.6 and 13.15, and an assistant will close that gap fluently, confidently, and without any access to the evidence, because closing gaps fluently is what it does.

That is also, and not coincidentally, where the assistant is genuinely useful. Language is the part of this work most analysts find hardest and slowest, and an assistant that produces eight candidate headlines in fifteen seconds is offering something real. The division of labor this section installs is therefore narrower than in previous chapters rather than broader: the assistant may generate candidates, and the analyst must verify every one against the view it sits on, one sentence at a time, before any of them reaches a canvas.

The failure modes are specific, and the first is the chapter's signature.

The overclaiming title is the modal failure and the most consequential. Asked to title a view, an assistant will reliably produce a sentence that is more interesting than the view, because more

interesting sentences are better writing and the assistant is optimizing for writing. The overclaim takes three recognizable shapes. It substitutes a causal or predictive verb for a descriptive one, so that "The four mature suburban and resort stores' four-week moving averages first reached the weekly fixed-cost threshold in weeks 43, 46, 51, and 57" becomes "the suburban format pays for itself within a year" — a claim about the future, about all such stores, made from four observations, using a word the view's measure does not support, and quietly dropping the store that took fifty-seven weeks. It generalizes the population, so that a statement about four mature stores becomes a statement about the format. And it imports a number the view does not contain, usually because the analyst pasted the chapter's findings into the prompt along with the view description and the assistant, quite reasonably, used them. The defense is mechanical and is the AI in Practice routine below: verify each sentence against the marks in the view directly beneath it, in the view's own units, with no other evidence admitted.

The confident recommendation is the second, and it is the same failure moved from the title bar into prose. Asked to draft a recommendation paragraph, an assistant will produce something decisive, because decisive prose is what the genre calls for and hedged prose reads badly. What it will not do, unasked, is carry the conditions — the interval, the threshold, the reversal criterion — and what it will frequently do is convert a conditional recommendation into an unconditional one during the rewrite, on the grounds of clarity. An analyst who hands over "sign five, hold three, revisit in eight weeks" and receives back "the data supports proceeding with the expansion" has watched the assistant delete the finding and keep the tone.

The disappeared uncertainty is the third and follows from the second. Asked to make a paragraph "clearer" or "punchier," an assistant removes the interval before it removes anything else, because the interval is the clause that complicates the sentence. This is worth predicting explicitly before any such request, and it is worth checking by diffing the two versions rather than reading the new one, since the missing clause is invisible in a text that reads well.

The layout without a hierarchy is the fourth. Asked to propose a dashboard layout, an assistant produces a plausible arrangement — often a neat grid, sometimes with a row of headline numbers along the top — that reflects conventional dashboard appearance rather than any judgment about which view the decision turns on, since it has no access to that judgment unless the prompt supplies it. The output is the first draft of Section 13.1 with better spacing. The fix is to supply the KPI hierarchy of Section 13.4 in the prompt as an input and ask the assistant to implement it, at which point the task becomes mechanical and the assistant does it well.

The generic audience is the fifth. Prompts that do not specify the reader produce artifacts for a composite executive who does not exist — some finance vocabulary, some marketing vocabulary, a level of detail calibrated to nobody. Pasting the audience analysis of Section 13.2 into the prompt changes the output substantially and costs one paragraph.

The invented finding is the sixth and the most dangerous in this chapter specifically, because it arrives dressed as help. Asked to write dashboard copy from a set of view descriptions, an assistant will occasionally supply a connective sentence that asserts a relationship between two views — "the higher order value explains the faster coverage" — that no view establishes and

that the analysis may in fact contradict. It reads as synthesis. It is confabulation, and it is caught only by the rule that every sentence must be traceable to a specific view.

The accessibility and provenance omissions are the seventh and eighth, and they are omissions rather than errors: nothing in a request for a headline or a layout will produce a note about the filter state, a statement of grain, alternative text for a screen reader, or a check on contrast, and none of these will be volunteered.

Where the assistant earns its place, the list is genuinely long. Generating a dozen candidate assertion titles for one view, so that the analyst chooses rather than composes, is the highest-value use in the chapter and inverts the usual difficulty of the task. Rewriting a paragraph for a stated reader — "rewrite this for a chief executive who has not read the analysis and will spend ninety seconds on it" — is a task assistants perform well and analysts perform slowly. Producing the calculated-field syntax for a dynamic title that tracks a parameter, or the Prep expression that reshapes four bound columns into the band-outline vertices Lab 13.1 Part D needs, is boilerplate — though the vertex order must still be checked by eye, because an assistant will produce a plausible ordering that does not close. Drafting alternative text for each view, which accessibility increasingly requires and almost nobody writes, is a natural fit. Converting a statistical statement into decision units — the store-equivalents arithmetic of Section 13.9 — is arithmetic the assistant will do correctly if given the inputs and will invent if not, so the inputs must be given. And the adversarial use remains the best one: handed a finished dashboard's title, subtitle, and view list, and asked what a reader could reasonably conclude that the underlying views do not support, an assistant produces a list that will contain at least one item the author did not see, in under a minute.

The governing instrument extends the audit lineage a final time. Table 8.5's three transferable points and Table 12.5's three visualization points both continue to apply — a communication artifact is built on views, and a view that fails V1 does not become acceptable because its title is good. Table 13.6 adds what communication makes necessary, and its three points are ordered by how often they are violated.

Table 13.6

The Communication Supplement to the Audit (run with Tables 8.5 and 12.5)

Supplemental point	The check	Fails when
C1. Every sentence is earned	For each title, subtitle, annotation, and line of recommendation prose, name the view that supports it and confirm the claim is verifiable from that view's marks in that view's units; confirm the verb is one the evidence hierarchy of Section 11.10 licenses; confirm every number in a sentence appears in a view or in the certified metrics file the view is built from	A title asserts a cause or a payback over a descriptive view; a headline generalizes four observations to a format; a number in the copy appears nowhere in the marks; a connective sentence asserts a relationship no view establishes

Supplemental point	The check	Fails when
C2. The reader's takeaway matches the intended one	State the intended takeaway in writing before testing; run the three-second test on a representative reader; compare	The reader names a different subject; the reader reports the artifact is "about the new stores doing badly"; the reader cannot state any takeaway
C3. What was removed is disclosed	List every simplification made for the audience — views cut, categories grouped into "other," windows shortened, filters defaulted, filter scope narrowed, qualifications moved into tooltips — and confirm each is either immaterial to the decision or stated on the canvas	The default filter state is restricted and unlabeled; a qualification that changes the recommendation lives only in a tooltip; a category that would change the ranking is inside "other"

AI IN PRACTICE

Draft the Titles, Then Prove Them

The chapter's verification theme as a five-step routine, run on every artifact whose language an assistant touched. Step one is the request, written to separate generation from judgment: "Here is a description of one view from a dashboard: it plots store-age weeks on the horizontal axis and the four-week moving average of weekly contribution on the vertical, one line per store, for the four suburban and resort stores open more than eighteen months, drawn on a complete store-week scaffold so that weeks with no orders appear as zero rather than being skipped, with a horizontal reference line at the \$1,769 weekly contribution that carries the \$92,000 annual fixed footprint. The first week in which each line reaches the reference is week 43, week 46, week 51, and week 57 respectively. Write twelve candidate titles as complete sentences.

Do not use any number that I have not given you. Do not assert a cause, a forecast, a payback period, or a claim about stores other than these four. For each candidate, state in a second line exactly which marks in the view a reader would have to look at to verify it." That last clause is the whole routine in miniature: it forces the assistant to produce its own audit trail, and a candidate whose verification line is vague is a candidate that cannot be verified. Step two is the prediction, made before reading the candidates: write down how many of the twelve you expect to overclaim, and in which of the three shapes — causal or predictive verb, generalized population, imported number.

Step three is the audit: take each candidate, cover the description you supplied, look only at the built view, and decide whether a reader could confirm the sentence from the marks; reject on the first failure rather than repairing, because repairing a candidate reintroduces the analyst's own knowledge into a sentence that must stand without it. Step four is the grade: compare the overclaim count and shapes against your prediction and record both misses and false alarms, since a systematic error in your model of the assistant is more useful than any single caught title. Step five is the reader, and it is the step that does not exist in Chapter 12: show the surviving title on its view to someone who has not seen either, ask what it says, and confirm that her paraphrase is the claim you intended rather than the stronger one next door.

File the exchange per Appendix D. And carry the rule that gives the routine its force: an assistant can write a sentence, but only the analyst can sign it, and the sentence at the top of a chart in an executive meeting is signed whether or not anyone thought about it that way.

13.12 Hands-On Application in Tableau

The preceding sections specified a reader, chose an artifact, built a hierarchy, laid it out, gave it sentences, emphasized and subtracted, bounded its interactivity, and stated its uncertainty. This section builds it. Lab 13.1 constructs the go/no-go dashboard in four parts: the decision brief, the certified inputs, and the KPI specification written before anything is opened; the building of three views, two of them rebuilt rather than reused from Chapter 12 and one with no ancestor at all; the assembly of the canvas with its containers, its one action, and its declared fixed state; and the uncertainty view that carries the fiscal-year outlook into the room. Lab 13.2 tests it: a three-second test and a ninety-second task on a peer, followed by a structured critique exchange, a revision, and the one-page recommendation memo.

The convention of Chapter 12 continues. A Tableau lab's equivalent of a code cell is a build path, so the boxed element is labeled **Build** and is followed by the same plain-language account of input, transformation, and output. AI assistants may draft any build path or any sentence in this section, and every output — including every sentence — is predicted before it is produced.

13.12.1 The Chapter's Inputs and Where They Come From

This chapter uses more declared quantities than any other in Part III, and it uses them in assertion titles, which is the one place a number must never arrive without a derivation. Before any build path, therefore, the inputs are named, sorted into the two kinds they belong to, scoped to the population they describe, and reconciled.

Start with the population, because everything else inherits it and because getting it wrong is the error that survives every other correction. The decision on December 2 is eight sites under letter of intent, and all eight are **suburban and resort** formats. StyleCraft's fifth wave also contains flagship sites, and they are deliberately not in this meeting: they are not yet at letter of intent, they carry a different cost structure, and finance has supplied no footprint figure for them. They are a separate decision on a separate benchmark, and saying so in one sentence in the source band is what keeps this dashboard's evidence from being read as evidence about them. Every *store-level* figure on the screen therefore describes the same population — the footprint, the margin, the predecessor stores, the Wave 4 cohort, and the prospective sites are all suburban and resort — and the two chain-level figures that remain, the fiscal-year outlook and the chain averages the supporting views draw as references, are labeled as chain wherever they appear. What is true without qualification is that **nothing on this screen describes the flagship format**. The single most useful check a student can run on her own artifact is to write the population beside each number and confirm that every entry reads either "suburban and resort" or "chain," and that no entry reads both.

One definition has to come before both kinds, because every figure on the screen sits downstream of it. **Weekly contribution** in this chapter is `line_revenue` less `line_cost`, summed within a store and a store-age week, and it carries weight only if `line_cost` is what the name implies. In the certified extract it is: `line_cost` carries all variable product and fulfillment cost assigned to the order line — merchandise cost, inbound freight, payment processing, and pick-and-ship — so weekly contribution is what remains after the costs that vary with the sale and before the fixed footprint the threshold represents. That is the property the whole dashboard depends on, because the comparison it draws is between contribution and rent, staffing, and allocated overhead, and a quantity that had not yet absorbed variable fulfillment cost would be compared against a footprint it cannot honestly be compared against.

Check this rather than inheriting it, and say what you find. If your extract's `line_cost` carries merchandise cost alone, the quantity you have computed is **gross profit**, not contribution: relabel it throughout, and either subtract the remaining variable costs before the comparison or state plainly that the threshold is being compared against a figure that has not yet absorbed them. The distinction changes nothing visually and changes what every title on the screen is entitled to say, which makes it exactly the kind of definition that gets inherited silently and should not be. The same applies to the 46 percent: it is a contribution margin only because the numerator is contribution, and it would be a gross margin if the numerator were gross profit.

The first kind of input is a **finance planning assumption**: a figure supplied by a person, for a purpose, that the extract cannot produce and does not contain. There is exactly one, and it is the \$92,000 annual fixed footprint — rent, staffing, and allocated overhead for a new StyleCraft store in the suburban and resort formats, which finance prices together. It is a forward-looking planning number carrying the finance partner's judgment about lease and staffing costs in the markets under letter of intent, and it is used here as declared. Two consequences follow and both belong on the canvas. It must be attributed to finance rather than to the data, and it must be declared before the results are inspected, because a threshold chosen after the crossing weeks are known is not a benchmark but a description — which is the reference-line discipline of Section 12.9 applied to the one number the whole dashboard turns on.

The second kind is a **certified derived quantity**: a figure computed from `orders_certified` by a stated calculation, published in the companion repository, and reproducible by any student who runs the calculation. Four files carry them.

- `chapter13_store_week_scaffold.hyper` — one row per physical store per store-age week, carrying weekly revenue, weekly contribution, the four-week moving average of contribution, and the predecessor range's lower and upper bound at that week; with the band-outline rows of Part A step 3 unioned in. This is the extension to contribution of the scaffold Lab 12.1 Part C built.
- `chapter13_dashboard_metrics.csv` — the scalar figures of Table 13.7, together with each store's first observed and first four-week-average coverage week.
- `chapter13_fy27_outlook.csv` — the fiscal-year contribution outlook at month grain, its interval bounds, and its certified horizon totals.

- `chapter13_prospective_sites.csv` — the eight sites under letter of intent, each with its proposed format, its trade-area population-density band, its median-household-income band, its matched existing comparable, and the flags recording eligibility and match.

The outlook file needs a paragraph of its own, because it is not a file Chapter 10 produced and because two decisions were made in building it that a reader is entitled to see.

The first decision is the **measure**. Chapter 10 forecast chain revenue. This chapter forecasts chain **contribution**, by rerunning Chapter 10's time-ordered model-comparison procedure on the daily contribution series and using the best-performing contribution model. Note the wording: not *the same model refit*, but *the same selection procedure rerun*. The specification that won on revenue may not win on contribution, whose variance structure, seasonal amplitude, and response to promotional weeks all differ, and assuming the winner transfers is the kind of inheritance this chapter spends a section warning against. The reason is arithmetic honesty rather than preference. The dashboard's uncertainty annotation converts an interval into store-equivalents, and the denominator of that conversion is a mature store's annual *contribution*. Converting a chain-level *revenue* interval into contribution would require multiplying it by a margin, and the only margin available is the 46 percent computed on the suburban and resort population — which is not the chain's margin, and assuming it is would be an unestablished claim buried inside an arithmetic step where no reader would find it. Forecasting contribution directly removes the step and the assumption together. The 46 percent margin still appears in this chapter, but only where it belongs: converting the weekly contribution threshold into the weekly *revenue* a suburban or resort store must take to reach it, which is a within-population conversion.

The second decision is the **horizon**, and it is where the audit of the previous draft was hardest to answer honestly. The certified transaction extract ends June 30, 2026, so the model's forecast origin is July 1, 2026 and its horizon runs to June 30, 2027. The meeting is December 2. Disclosing that five months of the horizon have already elapsed is better than concealing it, and it is not good enough: July through November are no longer uncertain future values, and presenting them inside a prediction interval two days before eight leases price would price risk that has already resolved.

The file therefore carries an **actual-to-date plus forecast-to-go** structure, and it is worth following the seam carefully, because the two halves come from different places. July through November are **actuals from finance's monthly close** — a separate certified source, at month grain, which the company has long since reported and which no analyst's transaction extract is needed to produce. December through June are the selected contribution model's forecast, from the July 1 origin, with the interval re-derived on the December-through-June sum from the same simulated paths. The fiscal-year figure is therefore the sum of a known quantity and an uncertain one, and only the second half has a range around it.

Two limitations follow and both belong on the canvas. The model has not seen July through November: it was fit on daily data ending June 30, and the analyst cannot refit it on months for which no transaction-level extract exists yet. The range on the remaining seven months is

therefore conditional on a fit that stops in June, which is why *refit when the extract refreshes* is one of the memo's reversal conditions rather than a housekeeping note. And the five elapsed months supply something free that the previous version threw away: the model predicted them, finance reported them, and comparing the two is a calibration check available at no cost. If the reported actuals sit near the middle of what the model expected for those months, the interval on December through June deserves more weight; if they sit outside it, the interval deserves less, and the memo should say which. The general lesson is that a forecast presented long after its origin is not merely stale — it overstates its own uncertainty about the part that has already happened, and it wastes the evidence that part contains.

One note on the status of everything below, because it governs how the table is used. Every figure in Table 13.7 is a **certified case value**: computed by a stated calculation from the certified extract, published in the companion repository with the answer key, and quoted in this chapter's titles, memo, and summary as the case's established result. None of them is provisional, and none is quoted with a hedge, because a chapter that hedged its own numbers could not teach assertion titles at all. What the student owes is reproduction, not adjustment. Recompute every row before building anything; where your figure differs from the published one, the difference is the finding and it goes in the audit — you do not move your number toward the table, and you do not move the table toward your number without saying so.

Two rows deserve a word about sensitivity, which is a different thing from provisionality. The coverage weeks and the Wave 4 classification both depend on the smoothing rule declared below, and a store sitting close to the range's edge at week 17 can be classified differently under a four-week average than under raw weekly contribution. That is a property of the measure rather than a weakness in the figure, and it is why the rule is declared before it is applied and why the reproduction step exists.

Table 13.7 reconciles every figure the dashboard displays and names the population each one describes.

Table 13.7

Reconciliation of the Figures Used on the Go/No-Go Dashboard

Figure	Population	Source and calculation	Verified result
Annual fixed footprint	New suburban and resort stores	Finance planning assumption, declared before results were inspected; not derived from any extract. No figure supplied for the flagship format	\$92,000 per store per year
Weekly fixed footprint	Same	$\$92,000 \div 52$	\$1,769.23

Figure	Population	Source and calculation	Verified result
Weekly contribution	Suburban and resort stores	orders_certified; $\Sigma(\text{line_revenue} - \text{line_cost})$ within store and store-age week, where line_cost carries all variable product and fulfillment cost	Revenue after variable cost, before the fixed footprint
Contribution margin	Suburban and resort stores	orders_certified; $\Sigma(\text{line_revenue} - \text{line_cost}) \div \Sigma(\text{line_revenue})$, on the same cost definition	46 percent
Weekly revenue at coverage	Same	$\$1,769.23 \div 0.46$	\$3,846.15
First observed coverage week	The four mature suburban and resort stores	Scaffold; first store-age week whose weekly contribution reaches \$1,769.23	Weeks 38, 42, 44, and 51
First four-week-average coverage week	Same	Scaffold; first store-age week whose four-week moving average of weekly contribution reaches \$1,769.23	Weeks 43, 46, 51, and 57
Mature-store annual contribution	Same	orders_certified; trailing 52-week contribution for each of the four stores over the year ending June 30, 2026, then averaged across the four. Not a weekly mean annualized, which would inherit whatever seasonality the sampled weeks happened to carry	Approximately \$137,000
Predecessor range	Same	Scaffold; minimum and maximum of the four-week average across the four stores at each store-age week 3–25, the average being null before week 3	23 bound pairs, one per week; see file

Figure	Population	Source and calculation	Verified result
Wave 4 position at week 17	The eight Wave 4 suburban and resort stores	Scaffold; count above, within, and below the range at store-age week 17, on the four-week average	2 above, 5 within, 1 below; the classification most sensitive to the smoothing rule
FY2027 contribution outlook	Chain	chapter13_fy27_outlook; July–November reported by finance, plus December–June forecast from the contribution model selected under Chapter 10's validation design, fit on certified daily history through June 30, 2026	\$1.25 million (\$0.51m reported + \$0.74m central estimate)
80 percent range, forecast-to-go	Chain, December–June only	Simulated paths summed over the seven remaining months, then 10th and 90th percentiles taken. Not reconstructible by summing the monthly bounds	\$0.62 million – \$0.86 million
FY2027 outlook range	Chain	Actual-to-date added to each end of the forecast-to-go range	\$1.13 million – \$1.37 million
Interval in store-equivalents	—	$(\$0.86m - \$0.62m) \div \$137,000$ end to end; half that around the estimate	1.8 stores end to end; ± 0.9 around it
Average order value	Chain; NYC Suburban	orders_certified (Chapter 12); $SUM(line_revenue) \div COUNTD(order_id)$	\$108 chain; \$156 suburban
Repeat-order share	Chain; suburban	orders_certified (Chapter 12); orders after the customer's first purchase date $\div$ distinct orders	47 percent chain; 24 percent suburban
Prospective sites	The eight sites under letter of intent	chapter13_prospective_sites; eligibility and matching rules of Part A step 6	All 8 eligible; 6 matched (5 to comparables that reached four-week-average coverage, 1 to Scottsdale); 2 unmatched

The labs use `orders_certified` and `lines_certified`, the two extracts produced by the Tableau Prep flow of Lab 12.1 Part A, which reconcile to \$4,882,247.60 in total revenue across 45,206 orders and 96,844 order lines; confirm those totals against the answer key in the companion repository before treating any of them as a target. Students who did not build the extracts should do so before beginning; nothing in this chapter re-derives them.

One definition remains, and it must be fixed before it is measured, because the phrase students reach for first is the wrong one and the rule they reach for second is too weak to carry a lease. The measure this dashboard is built on is not a payback period. Payback ordinarily means the time taken to recover an accumulated investment, and nothing on this screen accumulates anything: the view compares a store's weekly contribution against a week's share of its fixed cost.

DEFINITION

Weekly Coverage: Observed and Four-Week-Average

First observed weekly coverage is the earliest store-age week in which a store's weekly contribution reaches the declared weekly share of its annual fixed footprint. **First four-week-average coverage** is the earliest store-age week in which the four-week moving average of that contribution reaches the same threshold. Both are coverage milestones rather than payback periods: they say when a store carried its own weekly fixed cost, and say nothing about whether it recovered its build-out.

The distinction is not pedantic, and which one is made primary is a decision about what the artifact is for. A single promotional week, a holiday week, or one large corporate order can lift a store above the threshold once without establishing anything durable, and a lease is a multi-year commitment that should not be justified by a one-week milestone. This chapter therefore makes **four-week-average coverage the decision measure** and keeps observed coverage on the same view as a supporting milestone, drawn as a lighter marker, because the gap between the two is itself informative: a store whose two weeks are close crossed on strength its neighboring weeks shared, and a store where they are far apart crossed once on something the surrounding weeks did not repeat.

Be exact about what the four-week measure does and does not establish, because the language students reach for overstates it. It is **not** a claim that the store held coverage, stayed above the line, or established durable coverage. A trailing average can be lifted over a threshold by one strong week and drop below it the following week, and the milestone says only that at some point a four-week window averaged at or above the line. It is a smoothed first crossing, and it is preferred to the raw one because it is harder to produce by accident — not because it certifies persistence. A rule that did certify persistence would have to be declared as such: the fourth week of the first run of four consecutive raw weeks at or above the threshold, or the first week of a declared run of consecutive averaged weeks.

Either is defensible and neither is used here; what is not defensible is using the smoothed rule and describing it in the language of the stricter one.

Whichever rule is used, it must be declared before it is applied, and it must be computed rather than read off a plotted line. A line crosses a threshold somewhere between two weekly marks; the eye will pick the wrong one, and a value a student obtains by inspection is a value the next student cannot reproduce.

Source: Course concept developed for this guide, informed by the reference-line discipline of Section 12.9 and the baseline discipline of Section 2.7.

Note what the store-week scaffold does and does not do for these two measures, because the scaffold's value is easy to overstate in exactly one direction. Store age is derived from the opening date rather than from the row count, so a week absent from the transaction extract does not renumber the weeks after it: dropping a zero-order week before week 38 does not make week 38 arrive earlier, and **first observed coverage returns the same answer with or without the scaffold**. What the scaffold genuinely changes is threefold. It changes the *drawn path*, since Tableau otherwise connects across the gap and renders a fortnight of nothing as a straight line through it. It changes the *predecessor range*, because a minimum computed across four stores at each week is biased upward wherever one of those stores' zero weeks has been dropped, which narrows the region at its floor and can move a Wave 4 store from inside it to below it, or the reverse. And it is a **precondition** for four-week-average coverage, because a four-week moving average computed over an extract with holes in it averages the wrong four weeks. Making four-week-average coverage the primary measure is therefore also what gives the scaffold its work in the primary view rather than only in the comparison.

13.12.2 Lab 13.1, Part A: The Decision Brief, the Scaffold, and the KPI Specification

Almost nothing is drawn in this part. That is the point of it, and it is the part students skip and then rebuild the dashboard twice.

BUILD

The certified inputs and the written specification

Step 1. Extend the Chapter 12 scaffold to contribution. Open the Prep flow from Lab 12.1 Part A. On the branch that produced the store-week scaffold of Part C, add `Weekly Contribution` as the sum of `line_revenue` less the sum of `line_cost` within each store and store-age week, and set it to zero — not null — for any store-week in which the store was open and observable but recorded no orders. Leave it null where the store was not yet open or the window had not closed.

Add `Contribution 4wk Avg` as a Prep multiple-row calculation, specified completely rather than described: **group by** the store identifier, **order by** `Store Age Weeks` ascending, **compute using** `Weekly Contribution`, over a window of **the current row and the previous three**, and return null until four observable scaffold weeks exist for that store rather than averaging a short window (Tableau, 2026k). Every one of those five properties has to be stated, because a moving calculation with any of them left to a default computes a different number and the difference is invisible in the output. Extend the scaffold's store-age range from 0–25 to 0–64, since four-week-average coverage arrives later than observed coverage and the view must reach it. Output as `chapter13_store_week_scaffold`.

Step 2. On the scaffold, compute the predecessor range at each store-age week: the minimum and maximum of `Contribution 4wk Avg` across the four mature suburban and resort stores at that week, output as `Band Low` and `Band High`. Use the smoothed measure, not the raw weekly one, and use it because the primary view uses it: the decision rule is four-week-average coverage, and a comparison view that plotted raw weekly paths against a raw weekly range would evaluate the new stores on a measure the recommendation does not turn on. Identical axes do not make differently smoothed quantities comparable. These are per-week values, and note where they begin: the four-week average is null until a store has four observable scaffold weeks, so the range exists for store-age weeks **3 through 25** and not before.

That is 23 pairs, not 26, and the first three weeks of the axis carry no region at all — which is correct and should be left visible rather than patched, because a comparison band over a store's first fortnight would be built from a window that does not exist. Compute them in Prep, not in the view.

Step 3. Reshape those bounds into band-outline rows, which is what Tableau will actually draw. For each band, emit one row per vertex in drawing order: walking the lower edge left to right, then the upper edge right to left, so that the sequence traces the outline of the region and closes on itself. Four columns carry it — `series_type` set to `band`, `band_level` naming which band the vertex belongs to, `path_order` numbering the vertices from 1, and `band_value` holding the vertex's vertical position at that `store_age_week`. For the predecessor range over weeks 3–25 that is 23 vertices at `Band Low` from week 3 to week 25, followed by 23 vertices at `Band High` from week 25 back to week 3, or 46 rows.

Vertex 1 and vertex 46 share a horizontal position and not a vertical one — one sits on the lower edge and one on the upper — and Tableau closes the shape by drawing the segment between them, which is the left edge of the region. If a literally coincident final vertex is preferred, emit a 47th row repeating vertex 1. Union those rows into `chapter13_store_week_scaffold` beside the store rows, which carry `series_type` set to `store`. One extract, two kinds of row, one shared horizontal axis.

Step 4. On the scaffold, compute `First Observed Coverage Week` per store as the minimum store-age week whose `Weekly Contribution` reaches 1769.23, and `First 4wk Coverage Week` as the minimum store-age week whose `Contribution 4wk Avg` reaches the same value. Output both to `chapter13_dashboard_metrics` **and join them back onto every scaffold row for their store**, on the store identifier. Both are per-store scalars, and View 1 is built on the scaffold rather than on the metrics file, so a field that lives only in the metrics file cannot be dropped onto a shelf without a relationship the lab never establishes. Joining in Prep is the simpler of the two available answers; relating the metrics file to the scaffold on the store identifier is the other, and either must be chosen deliberately.

Reproduce every row of Table 13.7 and record any mismatch as a finding.

Step 5. Revalidate the outlook model on its new target before using it. Chapter 10 selected its model on the *revenue* series, and the model that wins on revenue is not necessarily the model that wins on contribution — a series with a different variance structure, different seasonality amplitude, and a different response to promotional weeks. Rerun Chapter 10's evaluation ladder on the daily contribution series without shortcuts: refit the naive and seasonal-naive baselines; refit each candidate under the same time-ordered rolling-origin design and the same holdout; compare errors in contribution dollars rather than in percentages; and either confirm that Chapter 10's specification still wins or adopt the contribution-series winner and say so. Publish the comparison table beside `chapter13_fy27_outlook`.

Until that ladder has been rerun, the outlook is a well-explained figure produced by a model validated for a different target, and no title should quote it.

Step 6. Build `chapter13_prospective_sites` from the eight letters of intent: one row per site, carrying its proposed format, its trade-area population-density band, its median-household-income band, its matched existing comparable, and two flags. Write both rules down before applying either.

Eligibility: a site is eligible for comparison against this dashboard's evidence only if its proposed format is suburban or resort; a flagship site requires a separate footprint, a separate margin, and a separate predecessor range, none of which exist here. **Match:** among eligible sites, a site matches an existing store when all three of the following hold — the same format, the same trade-area population-density band, and the same median-household-income band.

Tie-break: where more than one existing store qualifies, take the one with the smallest standardized distance across density and income, and record every candidate alongside the selection, because "closest comparable" is not reproducible unless the distance and the tie-break are both written down. The conjunction is deliberate. Format alone would match every eligible site to something, since every eligible site is suburban or resort and so is every existing comparable, and a rule that matches everything classifies nothing. Record the resulting counts. The memo's recommendation is a rule before it is a count, and a count without a rule is an opinion with a number attached.

Step 7. Write the decision in one sentence, naming the decision-maker, the choice, and the deadline. The form is: *StyleCraft's chief executive must decide by Friday whether to execute eight suburban and resort leases in full, in part, or not at all*. If the sentence requires an "and," the artifact is being asked to support two decisions and one of them belongs elsewhere.

Step 8. Write the audience analysis of Section 13.2 as five lines: who reads it, what decision, how long, what she knows, what would change her mind. Keep it to five lines; the constraint is the exercise.

Step 9. Write the intended takeaway as one sentence — the sentence you want the chief executive to say after three seconds. Seal it. Lab 13.2 grades it, and a takeaway written after the test is not a prediction.

Step 10. Complete Table 13.8 for the artifact: for every view, its tier, the reader's question it answers, the measure and grain, and the reference it carries — pairing each row with the assertion title drafted for it per Table 13.3. Seven rows is the target and eight is the maximum.

Step 11. For each of the five views inherited from Chapter 12, write one sentence naming the view it derives from and stating what must change in the rebuild; "nothing must change" is not an available answer, because if it were, Section 12.2 would be wrong. For the two purpose-built views, write one sentence instead stating the question they answer that no Chapter 12 view addressed, and why the pre-read did not need it.

Step 12. List every view and every control from the pre-read that will *not* appear, with one line apiece on where it goes instead — the memo's appendix, a tooltip, the workbook you bring to the room, or nowhere.

Input is the certified order-grain extract and the Chapter 12 Prep flow; the transformation is the extension of the store-week scaffold to contribution and its moving average, its reshaping into band-outline rows, and the computation of the derived quantities upstream of any view; the output is one extract carrying store rows and band rows, a metrics file, a site file, and a written specification. Steps 1 through 3 are the ones students are most tempted to skip, and skipping them changes the answer rather than the effort. `orders_certified` contains no row for a week in which a store took no orders, so a view built directly on it draws a line straight across the missing week, and a four-week average computed on it averages whichever four rows happen to be adjacent rather than the four weeks the calendar contains.

Step 6 deserves one further note, because the two rules it asks for do different jobs and students routinely collapse them into one. Eligibility asks whether this dashboard's evidence *applies* to a site at all, and it is answered by format, because the footprint, the margin, and the predecessor range are all suburban and resort quantities. Matching asks, among sites the evidence applies to, whether a specific existing store is a fair comparable, and it is answered by format *and* trade-area characteristics together. A site can be eligible and unmatched — the evidence applies to it in principle, but no existing store resembles it closely enough to say what its ramp will look like — and two of the eight are exactly that. Two failure modes sit on either side of the rule and both are worth naming, because students hit one or the other. A disjunctive match rule matches everything and therefore discriminates nothing, which is where the counts in Table 13.7 would collapse to eight and zero. A rule that folded eligibility into matching would let a trade-area resemblance stand in for a format the evidence does not cover, matching a flagship site to a suburban comparable while a later step declared that forbidden. Keeping the two rules separate and the second conjunctive is what makes the five-and-three split reproducible.

Table 13.8

The Go/No-Go Dashboard Specification

View	Tier	Reader's question	Measure and grain	Reference carried	Derives from
1. Coverage path	Primary	Do these stores cover their own cost, and when?	Four-week moving average of weekly contribution; one line per mature suburban or resort store, one mark per scaffold store-week	Horizontal line at \$1,769.23; first four-week-average coverage week labeled on each store, first observed coverage week marked more faintly	New; the pre-read did not answer the finance partner's question
2. Wave 4 ramps	Primary	Are the new ones on track?	Four-week average of weekly contribution by store age; one line per Wave 4 store, eight in all	The predecessor range drawn as a polygon whose outline is emitted in Prep; shared vertical scale with View 1	Lab 12.1 Part C, store-age facet, collapsed from 17 panels to one frame
3. Order value	Supporting	What is different about these customers?	Average order value at order grain, by <code>home_metro</code> ; one mark per metro	Chain average at \$108.00, entered as a constant from the metrics file	The Section 12.10 grain demonstration, drawn as a view for the first time. Lab 12.1 View A is <i>not</i> the ancestor: it answered the distribution question at order grain against the chain median, and Table 12.1 named a bar of the mean as the wrong form for that purpose
4. Category mix	Supporting	Is it a mix story or a pricing story?	Category revenue share by metro; one mark per category within metro	None; shares are self-referencing	Lab 12.1 View E

View	Tier	Reader's question	Measure and grain	Reference carried	Derives from
5. FY outlook	Supporting	How much of this is a guess?	Monthly chain contribution, July 2026 – June 2027, actual through November and forecast thereafter, with 80 and 95 percent bounds at month grain on the forecast months only	A labeled horizon-total block carrying the outlook and the range still open; interval width annotated in store-equivalents, both framings	chapter13_fy27_outlook, re-encoded
6. Repeat share	Diagnostic	What are you not telling me?	Repeat-order share by store_type; one mark per store type	Chain share at 47 percent, entered as a constant	Lab 12.1 View B
7. The outlier	Diagnostic	Which sites would you not sign?	Four-week average of weekly contribution by store age for the one Wave 4 store below the range	The same band polygon as View 2	Lab 12.1 Part C, one panel extracted

VERIFICATION CHECK

Before you run: predict what the dashboard will fail at, and write down three predictions sealed with the takeaway from step 9. Predict which of the seven views will turn out to be unnecessary — one of them will, and identifying it after the build rather than before is the ordinary outcome, so the prediction is worth grading. Predict the reading order a fresh reader will actually follow, as a numbered list of the seven views, and note that your prediction and your intended order are two different objects; if they match perfectly you have almost certainly predicted your intention rather than the reader. Predict the first question the chief executive will ask after the three seconds, and check whether the artifact answers it anywhere.

After you run: reconcile every row of Table 13.7 against your own outputs, and write the population beside each figure — every one of them should read "suburban and resort," "chain," or "derived," and any figure whose population you cannot name is a figure that has not been checked. Confirm that \$92,000 divided by 52 is \$1,769.23 and that this divided by 0.46 is \$3,846.15. Confirm that the contribution margin you compute on the suburban and resort store population rounds to 46 percent, and report the unrounded figure. Confirm that the scaffold contains exactly one row per physical store per store-age week in range, that zero-order weeks carry zero rather than null where the store was open, and that no store carries a non-null contribution for a week before it opened.

Confirm the band rows are 46 for the predecessor range, that `path_order` runs 1 to 46 without gaps, that the first and last vertices sit at store-age week 3 — at different heights, one on each edge, with Tableau's closing segment joining them, since an outline whose ends do not meet at a common horizontal position will not fill cleanly — and that no band row exists for weeks 0 through 2, which is the visible consequence of the four-week window's null rule.

Then run the three comparisons that show what the scaffold is actually for. First, recompute `First Observed Coverage Week on orders_certified` without the scaffold and confirm it returns the *same* four weeks, since store age comes from the opening date and dropping a week does not renumber the ones after it — a check designed to fail to find a difference, and worth running for that reason. Second, recompute `First 4wk Coverage Week` both ways and report whether it moved, and by how many weeks, for each store; the four-week window cannot be computed correctly on an extract with holes in it. Third, recompute `Band Low` at every store-age week without the scaffold and report the largest upward shift in the region's floor.

Finally, report the gap between each store's observed and four-week-average coverage week and say in one sentence what a large gap would mean about that store's early trading.

Investigate if: any store's observed coverage week is week 0 or week 1, which usually means the threshold was applied to revenue rather than to contribution; any store's four-week-average coverage week is *earlier* than its observed week, which is arithmetically impossible and means the moving average was computed on the wrong window; the predecessor range is flat across store age, which means `Band Low` and `Band High` were computed at table scope rather than per week; the site file records a match for an ineligible site, which means the two rules of step 6 have been collapsed into one; the contribution margin differs from 46 percent by more than a point, in which case the weekly revenue threshold changes and Table 13.7 must be rebuilt before any title quotes it; or `line_cost` turns out to carry merchandise cost alone, in which case every use of the word *contribution* in this chapter is wrong and must be corrected before anything is drawn.

13.12.3 Lab 13.1, Part B: Rebuild, Do Not Reuse

Chapter 12's Section 12.2 asserted that an analytical view should never drift into a communication artifact without being rebuilt. This part is where the assertion is cashed, and the instruction is literal: open new worksheets. Duplicating the pre-read's sheets and editing them preserves every default the analytical version was entitled to and the explanatory version is not, and the ninety seconds saved is reliably spent three times over.

BUILD

View 1 — the coverage path (primary)

1. New worksheet on `chapter13_store_week_scaffold`, not on `orders_certified`. The scaffold is the source for every view that plots a weekly path, for the reason Part A gave.
 2. Drag `Store Age Weeks` to **Columns** (continuous). Drag `Contribution 4wk Avg` to **Rows** with aggregation `Min` — the scaffold is already at store-week grain, so the aggregate returns the single scaffold value. Set the Marks card to `Line`, and `store_name` to **Detail**.
 3. Filter to the four mature suburban and resort stores. Filter `Store Age Weeks` to 0–64 so the frame reaches the latest four-week-average coverage week, and record that choice in the caption.
 4. **Analytics** pane → **Constant Line** on the vertical axis at 1769.23. Edit the label to read "Weekly contribution carrying the \$92,000 annual footprint (finance assumption, suburban and resort)" rather than the default value. A constant is correct here and a computed reference would be wrong: the threshold is external and was declared before the results were inspected, per Section 12.9 and Tableau (2026a).
 5. Color: place `store_name` on **Color** and assign one accent hue to all four, not four different hues. The reader's comparison is between the lines and the reference, not among the lines.
 6. Remove the legend and add direct labels: **Label** → **Line Ends** → **Label end of line**, showing `store_name`.
 7. Use `First 4wk Coverage Week` — joined onto every scaffold row in Part A step 4, which is why it is available here without a relationship — to place the annotations, rather than clicking the plotted line. Add it to **Detail**, then annotate the mark at each store's four-week-average coverage week: right-click that mark → **Annotate** → **Point** → label with the week number only ("week 43"). Four annotations, no more. Confirm that each annotated week is the first week at or above the reference on the moving-average line and not merely a week near it.
 8. Add the supporting milestone without letting it compete, and note that it needs a Marks card of its own — a field on **Detail** cannot change one point on a line into a different shape. Create `Observed Marker as IF [Store Age Weeks] = [First Observed Coverage Week] THEN 1769.23 END`, which returns a value at exactly one week per store and null everywhere else. Note the constant rather than the smoothed value: placing the marker on the *threshold line* at the right week says that this is where a single week first crossed, while placing it on the moving-average path at that week would invite the reader to take it for another measurement on the smoothed series, which it is not. Place it on **Rows** beside the line measure, right-click its axis → **Dual Axis**, → **Synchronize Axis**, and hide the duplicate header.
- Set that second Marks card to **Shape**, choose an open circle, reduce its size, and give it a lighter tint of the accent hue; separate measures receive separate Marks cards, which is what lets a line and a point shape coexist with different formatting (Tableau, 2026h). Then right-click the shape axis → **Move Marks to Front** so the markers are not lost under the lines (Tableau, 2026j). Leave them unlabeled, and edit the tooltip to say what the mark is — "first single week at or above the threshold" — so a reader who does hover is not left guessing. The gap between the hollow marker on the threshold and the annotated week on the curve should be readable without any of that.

9. Format: remove the worksheet border, remove the column gridlines, remove the axis title on the horizontal axis (the assertion title will say it), keep the vertical axis title with its units.

10. Title: replace the default with the assertion title from Table 13.3, and add a subtitle in smaller gray type stating the grain and both rules: "One line per store; four-week moving average of weekly contribution, on a complete store-week scaffold. Annotated week = the first week whose four-week average reached the threshold; hollow marker on the threshold line = the earlier single week that first crossed it."

Input is the store-week scaffold filtered to four stores; the transformation is none beyond the filter, because the moving average and both coverage weeks were computed upstream; the output is a four-line chart with one reference, four annotations, four supporting markers, and a sentence. This is one of the two views with no Chapter 12 ancestor — the pre-read never asked the finance partner's question, so there is nothing to rebuild — and it is therefore the cleanest available demonstration of what an explanatory view is when nothing is inherited. Seven of the ten steps are communication decisions rather than construction: one hue instead of four, direct labels instead of a legend, annotations instead of none, a restricted and justified window, a stripped frame, a second milestone deliberately subordinated rather than omitted, and a title that states a finding. The remaining three are not communication work at all; they are the difference between a number a reader can reproduce and a number the analyst read off a line.

BUILD

View 2 — the Wave 4 ramps against the predecessor range (primary)

1. New worksheet on `chapter13_store_week_scaffold`. Store Age Weeks on **Columns** (continuous). Filter Store Age Weeks to 0–25 — the full range the store lines occupy, and one week further would leave the region trailing off at an edge with nothing to close against. The band's vertices cover only weeks 3 through 25, for the reason Part A step 2 gave, so the first three weeks of the frame carry lines with no region behind them. Leave that gap; it is where the four-week average does not yet exist.

2. Build the region before the lines, because the region is the harder object and the lines are trivial. Its edges move week by week, which is what a Tableau reference band cannot produce: a band is defined between two values at table, pane, or cell scope and renders as a region of constant height across the pane (Tableau, 2026a). Nor can the two bound fields simply be dropped side by side on **Rows** and stacked — separate continuous measures on Rows produce separate axes, and stacking measures within one pane requires **Measure Values** and **Measure Names** (Tableau, 2026h), which then leaves every stacked segment sharing a single Marks card and therefore a single opacity, since opacity is a mark property of the card rather than of a measure (Tableau, 2026i).

The construction that avoids both problems is a **polygon**, drawn from the outline the Prep flow already emitted.

3. Create two calculated fields that split the unioned extract into its two kinds of row: `Band Value` as `IF [series_type] = "band" THEN [band_value] END`, and `Store Value` as `IF [series_type] = "store" THEN [Contribution 4wk Avg] END`. Note the measure: the four-week average, matching View 1 and matching the range's own construction. Nulls are dropped by each mark type, so the same extract feeds both layers without a blend.
4. Restrict the line layer to the intended cohort — and do it with a calculation rather than with an ordinary filter, because an ordinary filter on `Store Wave` or `store_type` would also discard the band-outline rows, whose store fields are null. Create `Cohort Filter`:

```
Cohort Filter =
  [series_type] = "band"
OR ( [series_type] = "store" AND [Store Wave] = "Wave 4" )
```

Place it on **Filters** and set it to True. The subtitle claims eight Wave 4 stores; this is the step that makes the claim true, and without it the frame contains every store in the chain while the sentence above it says otherwise.

5. Place `Band Value` on **Rows** as `MIN([Band Value])`. Set the Marks card type to **Polygon**. Drag `Path Order` to the **Path** shelf as a continuous dimension, and `Band Level` to **Detail**. The vertices now connect in the order Prep numbered them, tracing the lower edge left to right and the upper edge back again, and the outline closes on itself and fills.
6. Set the polygon's color to a pale tint of the neutral gray. One mark, one fill, no base to hide — which is the reason this construction is preferred to a stacked area, and the reason it extends without modification to the nested case in Part D.
7. Now add the lines. Place `Store Value` on **Rows** to the right of `Band Value`, right-click its axis → **Dual Axis**, and set that Marks card to Line with `store_name` on **Detail**. Right-click either axis → **Synchronize Axis**, then hide the duplicate header. Confirm the synchronization by reading one store's value against both axes before hiding.
8. Fix the layer order explicitly rather than relying on which axis is which. Right-click the line axis and choose **Move Marks to Front** (Tableau, 2026j), then confirm visually that every line is drawn over the region. Dual-axis layers can obscure one another, and this control exists precisely because the default is not something to depend on.
9. Set the vertical axis range to match View 1 exactly. This is the step that makes the adjacency of Section 13.5 work, and Tableau will not do it for you.
10. Color: one accent hue for all eight Wave 4 lines. Do not assign eight hues to eight stores.
11. Direct-label only Scottsdale, the one store tracking below the range. The other seven are the seven-of-eight the title refers to and are left unlabeled.
12. Title: the assertion title from Table 13.3. Subtitle: "Four-week average of weekly contribution. Shaded region = the range the four mature suburban and resort stores occupied at the same store age, week by week. Eight Wave 4 stores, all suburban or resort, 17 weeks old."

Input is the unioned extract of store rows and band-outline rows; the transformation is none inside Tableau beyond splitting the two row types and restricting the line layer; the output is a single frame carrying the comparison that Chapter 12 needed seventeen panels to make. Three things about this construction generalize and are worth naming. The first is the governing idea: **a region with a moving edge is a shape, not an aggregation**, so it is described upstream as an ordered set of vertices and Tableau is asked only to draw it. That is the same discipline as the floating-bar reference of Table 12.2 — a mark measured from an explicitly supplied reference rather than from zero — and it is why the chapter keeps computation in Prep. The second is that polygons compose: regions that do not overlap are simply separate sets of vertices distinguished by `Band Level`, which is exactly what the three edge-to-edge regions of Part D need. The third is step 4's lesson, which arrives whenever two kinds of row share one extract: **a filter written against a field that only one kind of row possesses will silently delete the other kind**, so the filter has to be written as a calculation that names both. A stacked **Measure Values** area with an opaque white base reaches a similar result and is the documented alternative if the polygon path proves awkward in your build; it costs the gridlines behind the region and cannot carry per-segment transparency.

Note also the trade being made, and make it consciously: the small multiples of Lab 12.1 Part C were the better analytical instrument, because they showed every store's individual path; this single frame is the better communicative instrument, because it shows one comparison and a ten-minute reader can only make one. Both are true, the second is not an improvement on the first, and an analyst who cannot say which object she is building will build the wrong one.

BUILD

View 3 — order value with the chain reference (supporting)

1. New worksheet on `orders_certified.home_metro` on **Rows**, and a calculated field `Average Order Value` as `SUM([line_revenue]) / COUNTD([order_id])` on **Columns**, formatted as currency with no decimals. Note what is being built and what is not. The pre-read's View A plotted the *distribution* of order values by metro, against the chain median, because the question then was whether suburban orders are uniformly larger or larger and more dispersed — and Table 12.1 was explicit that a bar of the mean is the wrong form for that question. The question here is different: the executive needs the level, once, against a chain benchmark. Same calculation as Section 12.10's grain demonstration, first time it has been drawn, and a form Chapter 12 would have rejected for Chapter 12's question.
2. Sort descending. Confirm the axis includes zero — bars encode by length, and Section 12.6's obligation does not relax because the audience is an executive.
3. **Analytics** → **Constant Line** at 108 → label "Chain average \$108." Enter it as a constant taken from the certified metrics file rather than as an average computed from the four displayed marks, for the reason Lab 12.1 View B gave: the chain figure is computed on the pooled denominator and is not the unweighted mean of the metros.
4. Color: the suburban bar in the accent hue, all others gray. One emphasis, per Section 13.7.
5. Direct-label the bar values; remove the value axis, which now duplicates the labels.

6. Title: the assertion title from Table 13.3. Subtitle: "One bar per metro; average order value computed at order grain; full certified window."

Input is the order-grain extract; the transformation is a ratio of two aggregations by metro; the output is a four-bar chart carrying one emphasized comparison against one drawn reference. Step 6's subtitle is not a formality: this is the number Section 12.10 showed can be computed three ways, the dashboard is asserting the order-grain version over the full certified window, and a reader who later reconciles it against a colleague's \$50.41 needs the grain on the canvas to resolve the discrepancy in the analyst's favor.

VERIFICATION CHECK

Before you run: predict what will have to be removed from each view's Chapter 12 ancestor, as a list, and grade the list afterward against what you actually removed.

After you run: run C1 from Table 13.6 on each of the three titles — cover your own knowledge, look only at the marks, and confirm that a reader could verify the sentence. For View 1, confirm the annotated weeks sit on the moving-average line rather than on a raw weekly value; confirm each hollow marker sits on the threshold line rather than on the curve, at its store's first single-week crossing; and confirm the markers are visibly subordinate to the annotations. For View 2, count the lines in the frame and confirm the count is eight and not seventeen; then confirm that "seven of the eight" and "within or above" are both things a reader could establish from the marks rather than things she must take on trust.

Verify the region is not flat: select three widely separated store-age weeks — say weeks 4, 12, and 24 — and confirm from the tooltip that the region's upper and lower vertices differ at each and are not the same pair repeated. Confirm the region begins at week 3 and not at week 0, and that the gap at the left of the frame is the three weeks in which no four-week average exists rather than a rendering fault; confirm too that it closes cleanly at both ends rather than trailing off. Finally, place Views 1 and 2 side by side at the size they will occupy on the dashboard and confirm that their vertical axes are identical.

Investigate if: the region renders as a rectangle of constant height, which means a reference band was used instead of the polygon; the region renders as a scribble or a bow tie, which means `Path Order` is on Detail rather than on the **Path** shelf, or is being treated as a discrete dimension; the region begins at week 0, which means the four-week average was computed over a short opening window instead of returning null; the region disappears when the cohort filter is applied, which means the filter was written against a store field rather than as the calculation of step 4; the lines are drawn under the region, which means **Move Marks to Front** has not been applied to the line axis; or the two axes are not synchronized, in which case every visual comparison in the frame is meaningless and no reader will be able to tell.

13.12.4 Lab 13.1, Part C: Assembly, Emphasis, and a Declared Fixed State

The views exist. This part makes them one artifact, and it is where the hierarchy of Section 13.4 either becomes visible or stays in the analyst's head. It is also where this chapter's most

counterintuitive design decision gets made, which is that the finished executive artifact carries no filter control at all.

BUILD

The canvas

1. **Dashboard** → **New Dashboard**. Set **Size** to Fixed, at the projection resolution of the room the meeting is in. Fixed is correct here because the consumption condition is known and the hierarchy is expressed in relative area; Automatic, Range, and device-specific layouts are the correct choices for artifacts published to unknown displays, and the choice should be recorded in the specification rather than defaulted (Tableau, 2026e).
2. Drag a **Vertical** container onto the empty canvas. Inside it, place four **Horizontal** containers: a title band, a primary band, a supporting band, and a diagnostic band. Build everything inside containers; floating objects will not hold their relationships when anything changes.
3. Title band: a **Text** object carrying the decision, not the data — "Wave 5: eight suburban and resort leases, decision Friday" — and, in smaller type, the date and the analyst's name. The person who signs the artifact should be on it.
4. Primary band: View 1 at roughly 40 percent of the canvas width and View 2 beside it at roughly 35 percent, axes aligned. This is the primary tier and it occupies the upper-left path.
5. Supporting band: Views 3, 4, and 5 as a half-height row beneath the primaries, each about a third of the width.
6. Diagnostic band: Views 6 and 7 at roughly half the height of the supporting row, with a lighter title treatment so their prominence matches their tier.
7. Add padding: select each container and set outer padding to at least 8 pixels, inner to 4. The whitespace is the grouping mechanism of Section 13.5 and the default of zero destroys it.
8. Add the source band: a Text object at the foot carrying the extract names, the certified window, the forecast origin, the populations, the state the screen is fixed in, and the refresh date. One band for the screen, not one per view. It should read something like: "Sources: orders_certified, chapter13_store_week_scaffold, chapter13_fy27_outlook, chapter13_prospective_sites. Physical stores only; online and app orders excluded. Coverage evidence covers suburban and resort formats only; flagship sites are a separate decision on a separate benchmark. Outlook: July–November reported by finance; December–June forecast from the contribution model selected under Chapter 10's validation design, fit on certified daily history through June 30, 2026. No filters applied; this screen shows the complete certified population."
9. Add no filter controls. This is a design decision and it needs to be made deliberately rather than by omission, so write the reasoning into the specification. The first draft carried seven. Five were removed because a ten-minute reader making a one-time decision was never going to use them and their visual presence made a false claim about what the screen was for. The remaining two were removed for a stronger reason, which is that no scope existed at which either was safe. A date-window control would move the marks on Views 3, 4, and 6 while leaving "\$156," "\$108," "24 percent," and "47 percent" standing in their titles and drawn in their reference lines, all four of which are full-window constants.

A `store_type` control is worse: applied to Views 1, 2, and 7 it can empty View 1 entirely, change the number of lines in a frame whose title says "seven of the eight," and — because the band-outline rows carry no store type — delete the predecessor region while leaving the sentence that describes it. Scoping a filter to selected worksheets narrows which sheets it touches; it does not keep a title true and it does not protect rows the filter has no value for (Tableau, 2026c). The honest number of filter controls on this artifact is zero, the state is instead declared in words in the source band, and all seven controls survive on a separate analytical tab in the workbook the analyst brings to the room.

10. Actions. **Dashboard** → **Actions** → **Add Action** → **Highlight**, source View 7, target View 2, running on hover, on `store_name`. The direction matters and is the point of the exercise. View 7 contains exactly one store, Scottsdale, and View 2 contains all eight, so every value the source can emit has a match in the target and the action fires for every possible selection. Reversed — source View 2, target View 7 — it would fire for one of the eight Wave 4 lines and do nothing, silently, for the other seven, because a highlight action operates on matching field values and a target that does not contain the selected value simply does not respond (Tableau, 2026b). View 1 is excluded from the action in either direction, since it contains only the four mature stores and shares no store name with the Wave 4 cohort at all.

Test the action by hovering Scottsdale's line in View 7 and confirming that its line separates from the other seven in View 2.

11. Parameter as navigation, with a title that tracks it. Create a string parameter `Metro Shown` listing the four metros and "All," defaulting to "All." On View 4, add a calculated field `Metro Filter reading` `[Metro Shown] = "All" OR [home_metro] = [Metro Shown]`, place it on Filters, and set it to True. Then create `Mix Title as IF [Metro Shown] = "All" THEN "Occasionwear is roughly three times the share of suburban revenue that it is of urban" ELSE "Occasionwear's share of revenue, " + [Metro Shown] + " only" END`, drag it to **Detail** — a calculated field cannot be inserted into a worksheet title unless it is present in the view — and then **Worksheet** → **Show Title** → **Insert** → `Mix Title` (Tableau, 2026f). Show the parameter control beside the view rather than anywhere that reads as a global setting.

12. Tooltips: on Views 1, 2, and 5, edit the tooltip to carry the store name or month, the store age, the weekly contribution or the month's bounds, and the grain sentence. Nothing load-bearing goes here, per Table 13.5.

13. Alt text: for each view, open the worksheet's **Accessibility** settings and edit the **alt text** so that a screen-reader user receives one sentence describing what the view shows and what its marks represent (Tableau, 2026d). Tableau's default alt text names the fields; replace it with the view's finding and grain.

Input is the seven worksheets; the transformation is spatial rather than numerical; the output is a single fixed-size screen whose geometry states the hierarchy and whose state is declared in prose rather than exposed in a control. Step 9 is the most consequential and the least obvious, and it is worth being clear about what it is and is not. It is not an argument that dashboards should not have filters; Table 13.5's first row still describes the conditions under which a filter communicates, and the analytical tab in the same workbook carries all seven. It is the application

of that row's own test to this specific artifact: a filter communicates when it narrows a population the reader already understands and cannot falsify anything on the screen, and on a canvas where six of seven titles quote fixed constants, no control on offer met the second condition. Step 11 is worth pausing on for the mirror-image reason, because it is the one interactive element that survived. In its default state View 4 carries an assertion title, which is a two-metro comparison; in any single-metro state that comparison is no longer on the screen, so the calculation returns a label title instead. That is legitimate precisely because the assertion has already landed in the default state and the parameter is answering a follow-up rather than delivering the finding. A parameterized view whose *default* state carries a label title has hidden its finding behind a control, which is the failure the last row of Table 13.5 describes.

VERIFICATION CHECK

Before you run: write down the state the screen is fixed in — population, window, forecast origin, exclusions — and predict which assertion titles would become false if a `store_type` control were added and set to Resort only. Name the four figures a date-window control would have falsified.

After you run: run three checks on the assembled canvas before anyone else sees it. The squint test — step back from the screen or blur the image until the text is unreadable, and note which object still dominates; if it is not View 1, the hierarchy is in your head rather than on the canvas, and the fix is size and position rather than a bigger title. The declared-state check from C3 of Table 13.6 — reload the dashboard fresh and confirm in writing that the source band states the population, the window, the forecast origin, the exclusions, and the fact that no filter is applied, because a screen with no controls on it makes no claim about its own completeness unless the words are there.

And the interaction inventory — list every interaction you built, and for each one name the reader's question it answers, state what the reader who never touches it will miss, and confirm by testing that it actually fires. Any interaction whose answer to the second question includes a finding is a design failure, and the finding moves onto the canvas.

Investigate if: the highlight action appears to do nothing, which means the source and target views do not share the field values; the parameter's title reads as an assertion in a single-metro state; the source band describes a state the screen is not actually in; or the alt text still names fields rather than describing the finding.

13.12.5 Lab 13.1, Part D: The Uncertainty View

This part builds View 5 separately because it is the view most likely to be built badly, and because it is where Chapter 10's work either survives the trip to the meeting or quietly does not. It reuses the polygon construction of Part B with three regions instead of one, and it adds the structure that makes a forecast presented five months after its planning year began still worth presenting.

BUILD

View 5 — the fiscal-year outlook with two nested interval regions

1. Understand the file's structure before drawing anything, because it is what the whole view is about. `chapter13_fy27_outlook` covers July 2026 through June 2027 and carries two kinds of month. July through November are **actual**, taken from finance's monthly close: contribution is known and there are no bounds. December through June are **forecast**, from a model origin of July 1, 2026, and carry `point`, `lo80`, `hi80`, `lo95`, and `hi95`, with the horizon rows holding the December-through-June totals quantiled from the simulated paths. The fiscal-year total is therefore a known quantity plus an uncertain one, and only the second half has a range around it. The file also carries, for each actual month, what the model predicted for it — which is what the calibration note in Section 13.12.1 asks the student to check.
2. Reshape the forecast months into band-outline rows exactly as Part A step 3 did, and union them with the monthly series into one extract. Three regions are emitted, and they are deliberately made **non-overlapping** so that no drawing-order problem can arise: `Lower_95` runs between `lo95` and `lo80`, `Central_80` between `lo80` and `hi80`, and `Upper_95` between `hi80` and `hi95`. Each carries `series_type = band`, its `band_level`, a `path_order` tracing its lower edge left to right and its upper edge back again, and `band_value` at each `Order Month`. The monthly rows carry `series_type = actual or forecast`, with `Order Month` and `value`. Every bound exists at month grain, which is the property the whole view depends on, and no band row exists before December, which is the property that makes the picture honest.
3. New worksheet. `Order Month` on **Columns** as a continuous date, spanning July 2026 through June 2027 and no further. Do not extend the axis past the fiscal year; a long empty tail shrinks the regions' apparent width, which is the axis-scaling question Exercise 13.17.7 returns to.
4. Do not generate a new forecast inside Tableau. This dashboard communicates an already verified result; replacing it with a second, unexplained model would break the analysis-to-communication chain, and Tableau's forecasting feature in any case displays only the single prediction interval currently configured, which cannot produce the two nested intervals this view requires (Tableau, 2026g).
5. Define the two measures the view needs, so that each mark type receives only the rows it should draw:

```
Band Value    = IF [series_type] = "band" THEN [band_value] END

Series Value = IF [series_type] = "actual"
                OR [series_type] = "forecast"
                THEN [value] END
```

6. Draw the regions. `MIN([Band Value])` on **Rows**, Marks card set to **Polygon**, Path Order on the **Path** shelf as a continuous dimension, Band Level on **Detail** and on **Color**. Assign the palest tint of the forecast hue to Lower 95 and Upper 95 and a visibly darker tint of the same hue to Central 80. Because the three regions are edge-to-edge rather than layered, the result reads as two nested intervals and nothing is hidden behind anything — which is the reason step 2 emitted three regions rather than two.

7. Add the series. Place `MIN([Series Value])` — the field defined in step 5 — on **Rows** as a second pill, right-click → **Dual Axis**, **Synchronize Axis**, hide the duplicate header, and set that Marks card to Line with `series_type` on **Detail** and on **Color**. Draw the actual months in the full forecast hue and the forecast months in a lighter tint of it, so that the change in certainty is visible in the line as well as in the regions. Note the consequence of putting `series_type` on Detail: it splits the series into two paths, so the line will break between November and December unless November is emitted a second time as an anchor row in the forecast segment.

Emit it, and emit the right value: the anchor repeats November's **reported actual** contribution, not the model's prediction for November, which the model produced from a July origin and which nobody is claiming as the month's outcome. Say why in the caption — the seam is communicated by the tint change and the seam line, not by a gap, and a gap invites the reader to think a month is missing. Do not attempt to draw one segment solid and the other dashed: a line's dash pattern is a mark property of the Marks card, set through **Path**, and it is not assignable to individual members of a dimension on that card (Tableau, 2026i).

8. Mark the seam. **Analytics** → **Constant Line** on the horizontal axis at December 1, 2026, labeled "Reported through November · forecast from here." This is the single most important non-data element on the view, because it is what tells the reader that the left half is history and the right half is a range. Note that this line marks where *reporting* ends, not where the model's origin sits: the model was fit through June 30 and has not seen the five actual months, and the subtitle has to say so.

9. Fix the layer order explicitly: right-click the line axis and choose **Move Marks to Front** (Tableau, 2026j), then confirm visually that the line is drawn over the regions.

10. Add the horizon total, which is the step most easily skipped and the one the assertion title depends on. The view plots months; the title states a fiscal-year figure, and no reader can verify a sum by looking at a line. Place a small text block inside the view's frame, in the upper left where it will be read before the marks, reading: "FY2027 chain contribution · \$0.51m reported (Jul–Nov) + \$0.74m forecast (Dec–Jun) = \$1.25m · FY range \$1.13m – \$1.37m, of which \$0.62m – \$0.86m is still open." Take every figure from the horizon rows of `chapter13_fy27_outlook`, which were produced by summing each simulated path over the seven remaining months and then taking percentiles of the totals.

Do **not** compute them in the view by summing the monthly bounds; adding up the edges of monthly intervals overstates the width, which is the aggregation rule Chapter 10 established and which this view is the natural place to violate.

11. Annotate once, at the right edge of the central region, with both store-equivalent framings: "The range still open on the remaining seven months spans about \$240,000 of contribution end to end — roughly 1.8 mature stores' annual contribution, or about ± 0.9 stores around the estimate."

12. Title: the assertion title from Table 13.3. Subtitle: "Monthly chain contribution. Left of the line, reported actuals from finance's monthly close. Right of it, the central estimate with the central 80% and 95% of simulated outcomes, from the contribution model selected under Chapter 10's validation design and fit on certified daily history through June 30, 2026."

13. Do not add a second measure axis, a growth-rate line, or a target line. This view has one job.

Input is the fiscal-year outlook file unioned with its band outlines; the transformation is none inside Tableau, because the three regions were described upstream as ordered vertices; the output is a continuous monthly series, unshaded over the reported months, with an uncertainty region that begins in December and widens across the forecast horizon. The series itself does not widen; the region around it does, and saying it the other way is the kind of shorthand that becomes a reader's misreading. Steps 8, 10, and 11 are the view's entire communicative payload, and it is worth noticing that all three are sentences or single marks rather than encodings: the seam line says where reporting ends, the horizon block makes the title verifiable, and the annotation makes the width meaningful. None of the three does another's job, and a view carrying only the regions has drawn an honest picture above a title no reader can check.

Step 1 also settles a question the previous version of this lab left open, and it is worth stating because it generalizes past this case. A forecast is not a fixed object with a fixed horizon; it is a statement made from an origin, and the distance between that origin and the room is part of the evidence. When months of a planning horizon have already happened, they are actuals and belong on the left of a line, not inside a band. Presenting them inside the band does not merely look stale — it overstates the analyst's own uncertainty, prices risk that has already resolved, and produces a range wider than the decision actually faces.

VERIFICATION CHECK

Before you run: predict what a reader will say the shaded region means, and write the prediction down; the two failure modes to watch for are "the contribution will be somewhere in there" stated as a certainty, and "they don't really know" stated as a dismissal, and the seam line and the annotation are what separate them.

After you run: confirm no shaded region appears to the left of the seam line, which is the check that the actual months carry no bounds. Confirm the regions widen: read the tooltip at December and at June and confirm the central region is thicker at June; a region of constant thickness means the bounds were not brought in at month grain. Confirm the three regions meet edge to edge with no gap and no overlap at every forecast month — `Lower 95` should end exactly where `Central 80` begins — which is a property of the file and a check on the reshape. Then confirm the horizon block's figures against the file: \$0.51 million actual plus a \$0.74 million central estimate is \$1.25 million, and the outlook range of \$1.13 million to \$1.37 million is the actual added to each end of \$0.62 million to \$0.86 million.

Then verify the store-equivalents arithmetic in the open rather than in your head. The width of the range still open is \$860,000 minus \$620,000, or \$240,000 of contribution; divided by the \$137,000 of annual contribution a mature store produces, that is 1.8 stores; halved, because the annotation quotes both, it is about ± 0.9 stores around the estimate. Notice what is *not* in that calculation: no margin. The forecast is already in contribution and the denominator is already in contribution, which is the whole reason the model was selected and fit on the contribution series rather than inherited from the revenue one.

Now compute two wrong answers deliberately — multiplying the \$240,000 by the 46 percent margin gives \$110,400, or 0.8 stores, and dividing the \$240,000 by the \$92,000 footprint gives 2.6 — and state in one sentence what each of them is actually measuring and why neither belongs on a canvas.

Investigate if: a shaded region appears in July through November, which means the actual months were given bounds; the region begins at the first month of the fiscal year rather than at the seam line; the regions render as a scribble, which means `Path Order` is not on the **Path** shelf; the line is drawn under the regions, which means **Move Marks to Front** has not been applied; the horizon block's figures do not match the file; or the summed edges of the monthly intervals were used as the range still open, which Chapter 10's aggregation rule forbids and which overstates the width.

Finally, remove the three regions and the seam line, look at the view, and write one sentence describing what the dashboard would have claimed without them. Keep that sentence; it is the exhibit for Section 13.15.

13.12.6 Lab 13.2, Part A: The Three-Second Test and the Ninety-Second Task

The artifact is finished. Nothing so far has produced any evidence that it communicates, and this part supplies the first of it. Work with a partner from the course who is building a different artifact; you will test hers afterward.

BUILD

The test protocol

1. Retrieve the sealed takeaway from Lab 13.1 Part A, step 9, and the sealed predictions from its Verification Check. Do not reread the dashboard first.
2. Write three further predictions: the first thing your reader's eye will land on, the first question she will ask, and the one element you expect her to misread.
3. Seat the reader in front of a blank screen. Display the dashboard for three seconds. Remove it. Ask one question: "What was that about?" Write her answer verbatim. Do not paraphrase it, and do not respond to it.
4. Display it again with one task: "You are the chief executive. Tell me what you would decide and why. Think aloud." Give her ninety seconds. Record, in this order, what she looks at first, second, and third; every question she asks; every hesitation; and her stated decision.
5. Say nothing for the entire ninety seconds. Do not explain any view. Do not answer any question — write it down instead and say you will come back to it. This is the hardest instruction in the chapter and the one that determines whether the test produces data.

6. Only after the ninety seconds are up, answer her questions and discuss.

Input is a finished artifact and an uncontaminated reader; the transformation is ninety-three seconds of structured observation; the output is a verbatim record and a set of graded predictions. The record is the deliverable, not the discussion afterward: what the reader said in the first three seconds is a measurement of the artifact's hierarchy, and it is worth more than her considered opinion at minute five, when she has become an expert on a dashboard she saw once. Note also what this check is and is not. It is a formative usability check on one reader, which finds gross hierarchy and language failures; it is not evidence that the artifact will work for a chief executive, and reporting it as though it were is the kind of overclaim Section 13.6 spent a page on.

VERIFICATION CHECK

Before you run: seal all six predictions in writing. A prediction written after the reader has spoken is not a prediction, and the entire value of the protocol is the gap between the two.

After you run: grade all six predictions explicitly, and write one paragraph on the pattern rather than on the individual misses. Three specific comparisons carry most of the information. Compare the reader's three-second takeaway to your sealed intended takeaway word by word: a mismatch in the *subject* of her sentence — she says the artifact is about new stores when you intended it to be about cost coverage — is a hierarchy failure and is repaired with size and position, while a mismatch in the *verb* — she says the stores are failing when you intended "on track" — is a title or emphasis failure and is repaired with language and color.

Compare her actual reading order to the order you predicted in Part A, and note every view she skipped entirely; a skipped view is not a view that failed to persuade, it is a view that was never read, and the two have different fixes. And compare the questions she asked against the questions Table 13.8 claims each view answers: a question the specification says is answered, asked anyway, means the view answers it in a way she could not find.

Investigate if: your reader restates your intended takeaway word for word, which is more often a sign that she has seen the artifact before or has been coached than a sign of a perfect design; she asks no questions at all, which usually means the task was not understood rather than that the artifact was; or every defect she surfaces is one you had already predicted, which means the prediction was written loosely enough to cover anything. Then state, in one sentence per defect, whether you will fix the artifact or accept the defect, and why — accepting a defect is a legitimate outcome, and recording it is what makes it legitimate.

13.12.7 Lab 13.2, Part B: The Critique Exchange, the Revision, and the Memo

The test produced defects. This part fixes them, and it adds the deliverable that carries the recommendation.

BUILD

The critique exchange

1. Swap artifacts with your partner. Each of you now critiques the other's, following the protocol of Section 13.10 in order and without skipping ahead.
2. Stage one, report: state what the artifact appears to say, in one sentence, before being told anything about it. Write it down before speaking.
3. Stage two, ask: ask what decision the artifact is meant to support and who reads it. Only now does the author speak, and only to answer these two questions.
4. Stage three, identify: name up to five specific elements that impede the reading, each as an observation about your own experience of the artifact rather than as a judgment — "I could not tell which of these two charts to read first," not "the layout is confusing."
5. Stage four, propose: for each identified element, propose one change. The author writes them down and does not defend, argue, or explain during any of the four stages.
6. Reverse roles and repeat.
7. Revise your dashboard against the combined evidence from Part A and the critique. For each change, record what evidence prompted it. For each piece of feedback you decline to act on, record why in one sentence.

BUILD

The recommendation memo

1. One page. Four parts, in this order: the recommendation, the evidence, the uncertainty, and the reversal conditions.
2. The recommendation is one sentence stating what should be done and to what scope, and it states the **rule** before any count. The conditional form of Section 13.9 is expected. It must also distinguish the two reasons a site is held, because they are different findings with different remedies: a site matched to a comparable that did *not* reach coverage has unfavorable evidence, while a site with no qualifying comparable has no evidence at all.

An imperative is not required — "We recommend signing the sites matched to comparables that reached coverage, and holding both the site matched to an unsuccessful comparable and the two with no qualifying comparable" is a recommendation, and so is the imperative form of the same sentence — but a sentence that states a finding rather than an action is not, and "the evidence supports the expansion" is a finding wearing a recommendation's clothes. A count may follow the rule; it may not replace it.
3. The evidence is one short paragraph per primary view, each stating the finding in decision units and naming the view it comes from.

4. The uncertainty is one paragraph containing the range still open in store-equivalents in both framings, the arithmetic that produced the conversion with its denominator named, the forecast origin and the share of the planning year already banked, the scope limitation — that every store-level figure on the screen describes the suburban and resort formats, that the outlook and the chain averages are labeled as chain, and that no footprint, margin, or predecessor range exists for the flagship sites in the pipeline — and the two facts on the diagnostic band: the 24 percent suburban repeat-order share against the chain's 47 percent, and the store below the predecessor range.
5. The reversal conditions are two or three observable events that would change the recommendation, each with a date or a threshold attached. At least one must name the refit that follows the next extract refresh, since the outlook's range rests on a model that has not seen the months already reported.
6. The appendix lists the views that came off the dashboard, with one line apiece on why, the finance assumption with its source named, the site-matching table from `chapter13_prospective_sites` showing how each of the eight sites was classified and by what rule, and the AI-Use documentation per Appendix D. Any count in the recommendation must be reconcilable to that table; a count the appendix cannot reproduce is an opinion with a number attached, and should be replaced by the rule itself.

Input is a tested dashboard and a critique record; the transformation is revision plus composition; the output is the pair of artifacts the December 2 meeting actually runs on. Note that the memo is written last and read first, which is the normal relationship between the two, and note that its first sentence is the sentence the whole chapter has been building toward — the one the analyst would say if she had thirty seconds, which Section 12.18.5 asked students to draft and promised would be rebuilt around. This is the rebuilding.

VERIFICATION CHECK

Before you run: state, in writing, how many sentences the artifact contains that make a claim — every title, subtitle, annotation, source-band line, and memo sentence. That count is the denominator for C1.

After you run: run the complete audit before either artifact leaves your hands — the three transferable points of Table 8.5, all three of Table 12.5 on every view, and all three of Table 13.6 on the artifact as a whole. C1 is run sentence by sentence, and the count is the deliverable: report how many sentences you checked and how many you changed. C2 is already run, in Part A, and the record of it is the evidence. C3 requires the disclosure list: every simplification, including each filter's scope, with a line stating whether it is material. Then run the final check, which is the one that matters most and takes the longest to accept. Hand the memo alone, without the dashboard, to a third person, and ask what decision it recommends and on what basis.

If she can state both correctly, the memo works and the dashboard is what makes it persuasive.

Investigate if: she cannot state the recommendation, which means the dashboard has been carrying an argument the analyst never actually wrote down — and an argument that does not exist independently of the picture will not survive the first person who reproduces the picture differently; she states a recommendation stronger than the one you wrote, which means the memo's conditions are buried behind its first sentence; or she asks what the interval means, which means the store-equivalent conversion did not travel out of the annotation and into the prose.

13.13 Marketing Interpretation and Managerial Insight

The labs produced a tested dashboard, a critique record, and a one-page memo. On December 2 the chief executive will look at the screen, form an impression, ask three or four questions, and decide. This section translates the artifact into the decision language it has to survive in, and — as in every chapter since Part II — it does so partly by exhibiting the wrong readings, because a communication artifact attracts a family of misinterpretation that differs in an important way from the ones Chapter 12 cataloged. Chapter 12's wrong readings came from two sources: charts that misencoded, like the truncated axis in the circulating deck, and charts that encoded honestly while inviting the wrong comparison, like the calendar-time facet that made the Wave 4 stores look weak. This chapter's come from a third place — charts that encode correctly, invite the right comparison, and are then asked to carry more weight than a picture can bear.

The first wrong reading is the chapter's signature and will arrive within ninety seconds of the screen going up. "This dashboard says we should sign the leases." It does not, and the reason is not that a dashboard is incapable of carrying a recommendation — this one carries a decision sentence in its title band, and a callout stating a decision rule would be a legitimate and often excellent addition. The reason is that the seven views themselves are a display: four are favorable to the expansion, two are unfavorable, and one is a range, and no arrangement of them constitutes the conditional sentence the evidence actually supports. What the room is doing is inferring a recommendation from emphasis, and emphasis is easily read as endorsement, particularly by an audience that wants a decision and will accept help forming one. The correction is not to weaken the design. It is to state the recommendation out loud, in the analyst's own voice, with its conditions attached, so that the room is responding to a claim somebody made rather than to an inference nobody is accountable for. The general form of the lesson is worth carrying into every meeting a student will ever present in: *charts and marks do not automatically constitute a durable recommendation, and when nobody states one, the audience infers one from the layout and attributes it to the analyst anyway*. Given that the attribution happens regardless, the only question is whether the analyst controls what is attributed to her — which is why a decision-support artifact pairs its evidence with explicit recommendation text and documented conditions.

The second wrong reading is the mirror of the first and is more expensive because it is quieter. "The analyst isn't sure, so this is really a judgment call for us." This is what an audience concludes when the uncertainty is presented and the recommendation is not, and it is the specific

failure the concept box in Section 13.9 was written against. Note what has happened when it occurs: the analyst has done the harder and more honest thing — she has shown the interval, named the weak store, put the 24 percent repeat share on the screen — and has been penalized for it, because the room has no way to distinguish *calibrated* from *unconfident*. The correction is structural rather than rhetorical. The four-part shape — recommendation, evidence, uncertainty, reversal conditions — puts the recommendation first specifically so that the uncertainty that follows is heard as qualification of a stated position rather than as the absence of one. An analyst who leads with the interval will be heard as having nothing to say, no matter how good the interval is.

The third wrong reading belongs to the diagnostic band and is the one that most often causes real damage in real meetings. "Scottsdale is underperforming — let's look at what's wrong there." The dashboard shows one Wave 4 store below the range its siblings occupy, which is a fact, and the room has converted it into a project. Two things are wrong with the conversion. The first is inferential. With eight new stores and only four predecessor paths, one visibly low store is not, by itself, evidence of a systematic operational failure — nor, and this is the part students get backwards, is it something the sample size *guarantees*. The predecessor range is the observed minimum-to-maximum spread of four stores rather than a statistical prediction interval, and how often a store should be expected to fall outside it depends on the underlying distribution, on the dependence between successive weeks, on how different the store formats are from one another, and on the precise rule used to classify a store as below. None of those are known here, so no probability can be quoted, and a claim that one low store was "close to certain" is exactly the kind of confident-sounding sentence C1 exists to catch. What can be said is the useful thing: Scottsdale should be treated as a diagnostic case and monitored against a rule declared in advance, rather than selected after inspection and subjected to an open-ended search for explanations — which is the post hoc hunting Section 11.8 named in a different costume. The second problem is decisional: the meeting's question is whether to sign eight *new* leases, and Scottsdale's performance is evidence about the *selection* of sites rather than about the operation of that store. The correction is to give the store its proper role, which is neither exhibit nor project: it is the case that defines the recommendation's condition. The recommendation is not "sign eight" or "sign seven"; it is "sign the sites matched to comparables that reached coverage, and hold both the site matched to Scottsdale and the two with no qualifying comparable until the ramp evidence separates" — at which point the outlier has done useful work and has not become an operational distraction that consumes a quarter.

A fourth reading is worth naming even though it is not exactly wrong, because it is the one that most often ends a meeting badly. "Can you add the online channel to this?" Some version of it will be asked, and the honest answer is usually yes, in about twenty minutes. The failure is to answer it in the room, live, in front of the dashboard, which is what the interactivity makes tempting and what the deliberate absence of any filter control on the screen was designed to prevent. Building a new view during a decision meeting produces an unaudited chart on a projector in front of the decision-maker, which is every mechanic in Chapter 12 waiting to

happen, and it converts a prepared argument into an improvisation. The professional response is to write the question down, answer it by end of day, and return the meeting to the decision — and the reason to name this in a textbook is that students taught to build fast are strongly inclined to demonstrate that they can.

The memo that accompanies the artifact has a fixed anatomy, and it is worth setting out in full because the exercises ask for it and because it is the shape of nearly every recommendation a marketing analyst will write for the rest of her career. It opens with the recommendation, which is a rule before it is a count: *sign the sites matched to comparables that reached four-week-average coverage; hold the site matched to an unsuccessful comparable and the two sites with no qualifying comparable; revisit all three in eight weeks, which the letters of intent make possible*. Applied to `chapter13_prospective_sites` under the two rules declared in Lab 13.1 Part A step 6, all eight sites are eligible — every one is suburban or resort — and six of them match an existing store: five to comparables that reached coverage, and one to Scottsdale, which has tracked below the predecessor range since week six. The remaining two match nothing, because no existing store shares their trade area's density and income bands. That distinction is not cosmetic — one site has unfavorable evidence and two have none, and the eight-week review resolves them by different means. The rule therefore signs five and holds three, and the count is stated only alongside the rule and the table that produces it. It deliberately appears nowhere in an assertion title, because no view on the dashboard displays the prospective sites and C1 of Table 13.6 does not make exceptions for the analyst's own recommendation. It states the evidence in decision units and by view: each of the four mature suburban and resort stores reached a four-week average at or above the weekly coverage threshold in weeks 43, 46, 51, and 57 respectively, which is the only claim the coverage view supports — four stores are four stores, no lease term or required ramp period has been established against which to call that range fast or slow, and the milestone says nothing about whether any of them stayed above the line afterward. The memo states all three limits rather than leaving them to be inferred; seven of the eight Wave 4 stores are within or above the range those four occupied at the same store age — two above it and five inside it — so the newest cohort's apparent weakness is consistent with store age rather than with performance; and the suburban customer's \$156 order against a chain average of \$108 is a basket-composition fact rather than a pricing one, which provides a testable expectation for markets with comparable customer profiles but does not guarantee that the pattern will reproduce. It states the uncertainty: the FY2027 outlook is \$1.25 million of chain contribution, of which \$0.51 million is already banked and \$0.74 million is a central estimate for the seven months from December, with the part still open running \$0.62 million to \$0.86 million — about \$240,000 end to end, worth roughly 1.8 mature stores' annual contribution or ± 0.9 stores around the estimate; the \$92,000 footprint is a finance planning assumption rather than a measured quantity, and every coverage figure moves with it; every store-level figure on the screen describes the suburban and resort formats and the two chain-level figures are labeled as such, while the pipeline's flagship sites carry no footprint, no margin, and no predecessor range, so nothing here bears on them; and the format's repeat-order share of 24 percent against the

chain's 47 percent is reflected in the contribution observed to date but remains a forward-looking acquisition-cost and durability risk that a second wave of stores would compound if it does not improve. And it states the reversal conditions: if the eight Wave 4 stores have not, by their store-age week 25 — the last week the predecessor range covers, and eight weeks from now, which is why the option period was negotiated to that length — held their position within or above that range, let the option lapse on the three held sites rather than executing; if the finance partner's footprint figure moves by more than ten percent when the leases price, every coverage week on the dashboard is recomputed before anything else is signed; when the transaction extract next refreshes, the model is refit on the months it has not yet seen and the outlook is rebuilt, and the recommendation is reopened if December lands outside the central 80 percent range for the month; and if suburban repeat-order share has not moved by the end of Q2, the expansion case becomes a customer-acquisition case rather than a real-estate one and belongs in a different meeting.

One property of that memo deserves a closing paragraph, because it is what makes this chapter's deliverable different in kind from Chapter 12's. Every artifact in Part II was valuable for the number inside it, and Chapter 12's pre-read was valuable for the errors it prevented. This chapter's artifact is valuable for a decision it makes *reversible*. The reversal conditions are the memo's most important sentences and the ones a junior analyst is most likely to omit, because they read like hedging and are the opposite: they are the mechanism by which a recommendation made under uncertainty on a Friday can be unmade in eight weeks without anybody having to admit to a mistake. Organizations that make decisions well are not organizations that are right more often. They are organizations in which being wrong is detectable and cheap, and the analyst who writes the detection criterion into the memo has done more for the quality of the decision than the one who got the point estimate closer.

13.14 Business Analytics in Practice

The chapter has practiced communication on StyleCraft. This section steps out of the fiction and into how visual communication actually functions inside organizations — where a single screen can acquire institutional weight nobody intended, where an entire company decided that the deck was the problem, and where a redesign produced measurable operational change that had nothing to do with aesthetics.

The first vignette is the daily flash and the governance burden that follows it. It is a composite drawn from recurring organizational practice rather than a report on one named firm, and the pattern recurs across retail, media, and consumer businesses: an analyst builds a simple morning dashboard — yesterday's revenue, orders, conversion, and a comparison to plan and to last year — for a specific manager who asked for it. It is useful. It circulates. Within a quarter it is the first thing the chief executive opens, and within two it is the number the company uses in conversation. What follows is a set of obligations that the original builder did not sign up for and that most analytics teams learn about the hard way. The flash must now be right by 7:00 a.m. every day, which makes it an operational system with an on-call expectation rather than a

workbook. Its definitions become the company's definitions, so that a change to how returns are netted or how a channel is classified is no longer a data decision but a communication event requiring notice, because executives who have watched a number every morning for a year detect a two-point discontinuity instantly and interpret it as a business event. Its availability becomes a dependency, so that a failed extract at 6:40 a.m. is a leadership problem. And its scope is under constant pressure, because every executive who reads it wants one more tile on it, and a screen designed around a three-second read degrades one reasonable request at a time into the undifferentiated grid this chapter opened with. The teams that handle this well institute two practices that connect directly to Part III. They version the definitions and announce changes in advance, which is Section 3.8's data dictionary promoted into a communication protocol. And they defend the hierarchy explicitly, treating a request to add a tile as a request to remove one, which is Section 13.7's discipline given organizational teeth. The practice lesson for a new analyst is that a widely read artifact is infrastructure, and that the moment to negotiate its ownership, its refresh guarantee, and its change process is before it becomes the chief executive's first screen rather than after.

The second vignette is the memo-first culture, and it is the clearest institutional expression of this chapter's argument that a screen alone does not carry a recommendation. Amazon is the best-documented example: the company's 2017 shareholder letter states plainly that it does not use PowerPoint in its senior meetings, that it uses narratively structured six-page memos instead, and that those meetings begin with everyone reading the memo silently (Bezos, 2018). The mechanism is worth understanding rather than admiring, because the practice spread to firms whose meetings it did not fit. What the memo does is force the argument into prose, where the connective tissue — *therefore, because, unless* — has to be written out and can be inspected. A deck lets a presenter place two charts on consecutive slides and let the audience supply the causal link, which is efficient when the link is sound and is the most common way an unsupported conclusion enters a decision when it is not. What the memo does not do is replace visual evidence; the memos contain charts, and appendices of them. The practice lesson for a marketing analyst is not that decks are bad. It is that if the recommendation cannot be written as a paragraph with its connectives intact, the analyst does not yet have a recommendation — she has a sequence of true statements and a hope that the audience will assemble them the way she would. That is exactly what Section 13.12.7's final verification check tests when it hands the memo to a third reader without the dashboard.

The third vignette is the redesign that produced a measured outcome, and it is worth telling because it is the vignette in which visualization work is evaluated the way the rest of the business is. It is likewise a composite. A campaign-performance dashboard exists, is opened daily by a dozen channel managers, and takes each of them roughly ten minutes to work through because the KPI hierarchy is flat and the pacing decision they actually make — whether to move budget between channels today — requires them to read four separate tiles and do arithmetic. A redesign restructures the same data around that decision: one primary view showing each channel's spend against its pacing target with the variance in dollars, the supporting detail one click away, and the

eleven tiles that nobody used removed entirely. The reported effects, where teams have bothered to measure them, are consistent in shape if not in magnitude — a substantial reduction in time-to-decision, a measurable increase in how often the pacing action is actually taken, and a decrease in the number of clarifying questions sent to the analytics team, which is the effect most visible to the analysts themselves. What did not change was the data, the model behind the pacing target, or the accuracy of anything. The practice lesson is that the return on communication work is real, is measurable by the operational behavior it changes rather than by anyone's opinion of the design, and is almost entirely invisible in the metrics analytics teams usually report about themselves. An analyst who wants credit for it has to instrument it, which means recording the decision the artifact is meant to support and checking, later, whether that decision got made faster or better.

In your first analyst job, these three compress into one expectation, and it is this chapter's closing thread: analytics careers compound through communication, because the last mile is the visible mile. The technical work of Chapters 4 through 11 is what makes an analyst employable, and it is also the part of the work that nobody outside the analytics team observes. What colleagues and managers observe is the artifact, the meeting, the memo, and the answer given to the awkward question — and their model of an analyst's competence is built almost entirely from those. This is not a complaint about superficiality; it is a description of how organizations necessarily allocate trust when they cannot inspect the work directly. The practical consequence is that the skills in this chapter compound in a way the technical ones do not. An analyst who is known to produce artifacts that can be read in three seconds and defended for thirty minutes gets invited to the meetings where decisions are made, which is where she learns what decisions actually require, which makes her next analysis better aimed. The analyst whose work is correct and unreadable is asked for numbers, and being asked for numbers is a career of answering other people's questions. The difference between the two paths is not intelligence and is not usually effort. It is whether the last two feet were treated as part of the job.

13.15 Ethics, Persuasion and Distortion

Section 13.14 ended on the observation that communication is what an analyst is judged by. This section takes up what follows from it, and it is the chapter's hardest material because it is the place where this guide's running ethics discussion — data use (Section 1.13), problem framing (Section 2.12), measurement design (Section 3.13), cleaning as editorial power (Section 4.14), honest summarization (Section 5.14), differential treatment of segments (Section 6.16), causal language (Section 7.15), acting on predictions about people (Section 8.15), fairness across groups (Section 9.15), forecast accountability (Section 10.15), experimentation on customers (Section 11.15), and misleading chart mechanics (Section 12.16) — arrives at a boundary that cannot be drawn by rule.

Chapter 12's ethics were, by comparison, easy. The five mechanics it cataloged are wrong in every context: a truncated bar axis misstates a proportion whether the analyst intended it or not, and the obligation it creates is to check. This chapter has no such list, because every technique in

it is legitimate and every one of them is also the mechanism by which a reader can be steered. Placing the cost-coverage view at the upper left in the largest frame is the hierarchy of Section 13.4 doing its job; it is also a decision that the reader will encounter the expansion's best evidence first. Coloring the four stores that reached the four-week-average coverage milestone in an accent hue and the rest in gray is the emphasis discipline of Section 13.7; it is also a decision about which stores the reader will look at. Writing "seven of the eight new stores are within or above the range their predecessors occupied" is an assertion title verified against its view; the same view honestly supports "one of the eight new stores has tracked below its predecessors since week six," and choosing the first sentence is choosing what the reader takes away. None of these misstates a quantity. All of them shape a conclusion. Hullman and Diakopoulos (2011) gave the general phenomenon a name and a taxonomy, cataloging the framing effects available in narrative visualization — what is included, what is omitted, how a view is annotated, what order the views arrive in — and observing that these rhetorical choices are unavoidable rather than optional, since every artifact makes them and an artifact that declined to make them would not be readable. That is what makes the chapter's ethics a matter of judgment rather than of checking, and it is why the boundary has to be argued rather than listed.

CONCEPT

Emphasis Directs Attention; Distortion Directs Belief

The working distinction this guide adopts is between techniques that change what a reader *looks at* and techniques that change what a reader *concludes is true of the data*. Emphasis, ordering, sizing, decluttering, and titling all direct attention: they make one true thing easier to find than another true thing, and a reader who spent an hour with the underlying evidence would arrive at the same understanding of the numbers, more slowly. Distortion changes the understanding itself: after encountering it, a reader believes something about the quantities that the data does not support, and would revise that belief on seeing the evidence. Chapter 12's mechanics are distortion by construction.

But distortion is available in this chapter too, without any mechanic, through three moves that are worth naming because they are the ones that occur. Selection: showing the four views that favor a conclusion and omitting the three that qualify it, so that the reader's picture of the evidence is not the analyst's. Framing: presenting a quantity in the form that produces the intended reaction — "repeat share improved 20 percent" for a move from 20 to 24 percent — where the alternative framing is equally true and produces a different one (Tversky & Kahneman, 1981). Aggregation: grouping the categories whose individual values would change the ranking into an "other" that hides them.

Each of these produces a false belief with true components, which is why the test for them cannot be *is every element accurate* and must instead be the counterfactual: *would this reader revise her conclusion if she saw everything I saw?* If the answer is yes, the artifact is distorting regardless of how carefully each of its elements was verified.

That counterfactual is the standard, and three practical obligations follow from it.

The first is the obligation of the omitted view, and it is the one this chapter's technique makes most acute. Section 13.7 requires decluttering, and decluttering requires removal; the go/no-go dashboard came down from the pre-read's eight views to seven, three of them cut and two purpose-built ones added, and from seven filter controls to none, and each removal was defended in terms of the reader's ten minutes. The obligation is to sort the removals into two piles. A view removed because it is redundant, or because the memo carries it better, or because it answers a question this reader will not ask, is a legitimate simplification. A view removed because it complicates the recommendation is not, and the fact that the same sentence — "it was cluttering the screen" — is available to describe both is precisely why the sorting must be done deliberately and in writing, which is what C3 of Table 13.6 exists to force. The practical test is the one the labs run: list every removal, and for each state whether the reader's decision would change if she saw it. Anything that survives that question with a yes goes back on the canvas, at whatever size, however inconvenient.

The second is the obligation of the disclosed state, and it is narrower and almost mechanical. An interactive artifact is in some state when the reader encounters it, and that state was chosen by the analyst. A filter defaulting to the last six months, a store type excluded because it "isn't comparable," a forecast whose origin is months before the meeting, a filter scoped to three worksheets rather than seven — each may be entirely correct and each is invisible to a reader who has no reason to check the controls. The standard is that the state must be stated in words on the canvas, not merely reflected in a control, because a control reading "Suburban, Resort" is a fact about an interface and a caption reading "Physical stores only; online orders excluded; suburban and resort formats only; July–November reported, December–June forecast from a model fit through June 30; no filters applied" is a fact about the evidence. The go/no-go dashboard makes this obligation sharper by carrying no controls at all, which removes the interface from the question and leaves nothing but the words: a screen with nothing to adjust says nothing whatever about what it contains unless the source band says it. This is the footer discipline of Section 12.9 carried into a medium where the artifact can change under the reader, and it is the reason the source band survived the decluttering that removed almost everything else.

The third is the obligation that comes with an audience that cannot check, and it is the one that grows over a career. Chapter 12's readers were colleagues with the extract and a week. This chapter's reader is a chief executive who will never see the file, and the readers of the artifacts these students will build in their first jobs are further away still — a board, a regulator, a customer, a journalist reproducing a chart out of a press release. Every step in that direction transfers responsibility from the reader to the analyst, and the transfer is complete when the artifact travels beyond the room where its author can be questioned. The working standard is therefore that a persuasive artifact must be built so that it would survive its own audience being fully informed: the analyst should be able to imagine handing the chief executive the entire pre-read, the certified extract, and an afternoon, and to expect that the chief executive would reach the same decision. An artifact that depends on the reader not having that afternoon is not a

communication artifact. It is an argument that requires an ignorant audience, and the fact that the audience is in fact busy rather than ignorant does not change what the artifact is doing.

One boundary closes the section and the guide's ethical thread. There is a version of this material that collapses into paralysis — every design choice shapes a conclusion, therefore every design choice is suspect, therefore the honest analyst makes no choices and produces the undifferentiated grid of Section 13.1. That conclusion is wrong and is worth rejecting explicitly, because it is where careful students land. A screen with no hierarchy is not neutral; it is a screen whose hierarchy was set by the order the views were built in, which serves nobody and which still produces a conclusion in the reader's head, just an accidental one. Refusing to choose is itself a choice with consequences, and it is a worse one, because it is unaccountable. The analyst's position is not that she should avoid influencing the decision — influencing the decision is the job, and it is why the organization employs her. It is that she should influence it with the evidence, in the direction the evidence points, having first checked that it points there, and having built the artifact so that a reader who went and looked would agree. That is a demanding standard and it is achievable, and the difference between meeting it and not is usually a written list of what was left out.

13.16 Chapter Summary

This chapter completed Part III's account of visualization by taking its second job, and it began by insisting that the second job is not a polished version of the first. An analytical artifact is built for a reader who has time, stake, and a reason to interrogate; a communication artifact is built for a reader who has a decision, a deadline, and no prior exposure to the evidence, and the design constraint is therefore the reader rather than the data. Audience analysis makes that constraint writable — who reads, what decision, how long, what she knows, what would change her mind — and the course's three-second reading test makes it measurable, by isolating what an artifact communicates before a cooperative reader decides to work at it. From there the chapter chose the object before designing it: exploratory artifacts let a reader investigate by her own path and explanatory ones deliver a finding along a path the author chose, and the go/no-go artifact turned out to be a data story delivered on a dashboard's canvas with a memo attached, which settled half a dozen later decisions including why its filter panel shrank from seven controls to none. The screen was then given a structure. The KPI hierarchy spatialized the metric discipline of Section 3.5 into primary, supporting, and diagnostic tiers, so that the two views the lease decision actually turns on — the store-age week in which a mature store's trailing four-week average of contribution first reached the \$1,769.23 weekly share of a declared \$92,000 annual footprint, and where the eight Wave 4 stores sit against the range their predecessors occupied — occupy the largest frames on the reader's entry path, and so that the case's two weakest facts sit on the canvas rather than in a backup slide. Layout converted the hierarchy into geometry through reading order, grouping, density, and whitespace, on the principle that position is a claim about importance whether or not the analyst intended to make one.

The chapter then supplied the techniques Chapter 12 deliberately withheld. A narrative arc organized the views around the finance partner's objection rather than around the strongest evidence, on the grounds that an audience holding an unaddressed doubt is not persuaded by anything else. The assertion title converted each view from a display of a quantity into a statement of a finding, at the cost of an obligation that a label title does not carry — that the sentence be verifiable from the marks beneath it, in the view's own units, with a verb the evidence hierarchy of Section 11.10 licenses and no number the view does not contain. Preattentive attributes supplied the emphasis toolkit, with the conjunction penalty and the conditions under which the effect weakens explaining why emphasis is a budget rather than a technique, and data-ink discipline supplied the subtraction that gives emphasis something to work against — a subtraction the chapter framed honestly as deleting one's own work and made survivable by moving rather than destroying. Interactivity was bounded to navigation and communication, with the rules that every finding must be visible in the default state, that a highlight action can only mark entities its target views actually contain and should therefore run from the smaller value set to the larger, that a filter's scope must be chosen and disclosed rather than applied to everything — and that a control comes off entirely when no safe scope exists, which is what happened to both of this artifact's surviving filters and is why the finished screen carries none, declaring its fixed state in the source band instead — and that a title over a parameterized view must be built to stay true in every state the control can produce. Uncertainty then received its own treatment, because Chapter 10's interval and Chapter 11's break-even had to survive a room that wants a number: stated as a central estimate with a central 80 percent range rather than as a most-likely value and a range that cannot be ruled out; structured as actual-to-date plus forecast-to-go, so that the five months of the planning year that had already happened by December 2 sat outside the band as history rather than inside it as risk; stated in contribution rather than revenue, so that the store-equivalent conversion needed no margin and carried no unestablished assumption; converted into decision units in both framings — 1.8 mature stores end to end, about ± 0.9 around the estimate; and drawn as regions whose outlines are described upstream as ordered vertices and drawn as polygons, because a reference band renders at constant height, a prediction interval does not, and two bound measures dropped side by side on Rows do not stack into a region on their own. The four-part shape — recommendation, evidence, uncertainty, reversal conditions — was the chapter's answer to the false choice between false confidence and useless hedging, and its fourth part is what turns an interval into a monitoring plan.

The last movement closed the loop from artifact to decision. A screen displays evidence and can carry a recommendation only as text somebody wrote, so it is paired with a one-page memo that states the recommendation, is written before the meeting because a recommendation that will not survive being written will not survive being questioned, and carries the reversal conditions that make the decision cheap to unmake. The usability check then supplied this chapter's verification theme, which is the first in the book that the analyst cannot run alone: predict what a fresh reader will say, show the artifact for three seconds, record the answer

verbatim, say nothing for ninety seconds more, and grade the predictions — with a mismatch in the subject of her sentence diagnosed as a hierarchy failure and a mismatch in the verb as a title or emphasis failure, and with the whole exercise reported as the formative check it is rather than as validation. The critique protocol made the same exchange survivable between peers by fixing its order. The AI section named the failure modes that follow from the assistant drafting the interpretation rather than the instrument — the overclaiming title in its three shapes, the confident recommendation that deletes the finding and keeps the tone, the uncertainty that disappears when a paragraph is made punchier, the layout that reflects convention rather than hierarchy — and installed the routine that requires every candidate sentence to arrive with its own verification line. And the ethics section drew the line Chapter 12 declined to draw: emphasis directs what a reader looks at, distortion directs what she concludes is true, and the test between them is a counterfactual rather than a checklist — would this reader revise her conclusion if she saw everything I saw? The obligations that follow are the sorted list of what was removed, the filter state and scope stated in words rather than shown in a control, and the standard that a persuasive artifact must be built to survive its own audience being fully informed. What the section explicitly refused is the paralysis that the argument invites, since a screen with no hierarchy is not a neutral screen but one whose hierarchy was set by accident.

StyleCraft's storyline closes here. The decision gets made on December 2 and the eight leases price on Friday, which ends a thread that has run since Chapter 2 arrived with a vague request to look at the expansion. It is worth noticing what carried it. Not the model — the forecast is one of seven views and the smallest claim on the screen. What carried it was a chain of small verified moves: a metric defined at the grain the question required, a baseline the result had to be read against, a threshold declared before the results were inspected, a scaffold that refused to let an absent week disappear, a coverage rule that asked for four weeks rather than one, an interval placed only around the part of the year that was still open, a sentence checked against the marks beneath it, and a reader who was asked what she saw instead of told what she should have seen. That chain is the guide's actual subject, and the dashboard is only the last link of it.

What remains is yours to build. Your instructor may release the Project #2 brief with this chapter, and its deliverable is the artifact this chapter teaches: an executive-facing dashboard or visual report, paired with a written recommendation, tested on a reader before it is submitted. Three appendices carry the instruments you will need and are worth opening before you start rather than after. Appendix D holds the AI-Use Documentation Template, which is what makes an AI-assisted artifact auditable rather than merely fast. Appendix E holds the Data Visualization Checklist, which is where Chapter 12's mechanics and this chapter's three communication points live as something you can run rather than remember. Appendix F holds the project rubrics, including Project #2's, which you should read before you open a canvas, because a rubric read afterward is a grading instrument and a rubric read beforehand is a specification.

And the standard is the one Chapter 1 set. AI may draft the chart, the title, the layout, and the paragraph; the student remains the analyst of record. Specify what the artifact must do and for whom, before building it. Predict what it will show and what a reader will take from it, then

verify both — the first against the marks, the second against a person. Explain every choice you made, including the ones the software made for you and you accepted. And document the exchange, so that the next reader of your work can see not only what you concluded but how you came to be entitled to conclude it. That workflow was introduced in Chapter 1 as a way to use a new set of tools responsibly. Thirteen chapters later it is simply what it means to do the work.

13.17 Exercises for Practice and Homework

The following exercises practice the chapter's main habits: state the reader before the design, name the decision the artifact must support, put the primary view where the eye lands first, write the sentence you want the reader to leave holding and then verify it against the marks, remove what does not help her believe it, show the uncertainty in units she cares about, and find out from an actual person what your artifact says. Not everything below is required. Core practice, which every student should complete: 13.17.1, 13.17.2, 13.17.3, and 13.17.4. Homework options, from which your instructor will assign a subset: 13.17.5 and 13.17.6. Instructor-led discussion, prepared for class rather than submitted: 13.17.7 and 13.17.8. Each exercise also carries its label individually. Your instructor may assign selected exercises as preparation for Project #2; consult the course schedule for deadlines. Exercise 13.17.5 produces the artifact the project is built on, and the Project #2 rubric in Appendix F should be read before starting it rather than after.

13.17.1 Concept Check (Required Practice)

Answer each in two or three sentences, in your own words.

- State the five questions an audience analysis must answer, and for each, name one design decision it determines.
- Explain what the three-second test measures and what it does not, state why this guide uses three seconds rather than the five used in first-impression testing generally, and state why the correct response to a failed test is to revise the artifact rather than the reader.
- Distinguish an exploratory artifact from an explanatory one, and give one example of an artifact built for an external audience that is nonetheless exploratory.
- Define a KPI hierarchy, and explain why a canvas of six equally sized views has communicated a claim the analyst does not believe.
- Explain what preattentive processing is, name two conditions under which the effect weakens, and state why a dashboard that emphasizes six elements has emphasized none.
- Distinguish data-ink from non-data-ink, and name two elements of non-data-ink that should nonetheless never be removed, with the reason for each.
- State the obligation an assertion title carries that a label title does not, and give one example of a title that is true but not earned by the view beneath it.
- Explain why a Tableau reference band cannot draw a prediction interval that widens with the horizon, why placing the lower and upper bound measures side by side on Rows does not produce a filled region between them, and what must be true of the underlying data before the interval can be drawn correctly.

- Define first observed weekly coverage, state why it is not a payback period, and explain why a complete store-week scaffold does not by itself move the first qualifying week later while it is indispensable to the four-week variant of the same rule.
- Explain the difference between emphasis and distortion as this chapter defines them, and state the counterfactual test that separates them.

13.17.2 Audience First (Required Practice)

For each of the following deliverables, write the five-line audience analysis of Section 13.2, then choose an artifact from Table 13.1 and justify the choice in one sentence, then state the intended three-second takeaway.

- A weekly summary of email campaign performance for four channel managers who will use it to decide where to move next week's sends.
- A one-time analysis, requested by a chief financial officer, of whether the loyalty program's discount cost is justified by the incremental revenue it produces.
- A monitoring artifact for a merchandising team that needs to notice, within a day, when a product category's sell-through falls behind plan.
- A summary of Chapter 11's A/B test result for a marketing leadership team deciding whether to roll the tested creative out to the full list.

Then answer the question the exercise is actually about. Two of the four are cases in which a dashboard is the wrong answer. Name them, state what the right artifact is, and explain what specifically about the reader or the decision rules the dashboard out.

13.17.3 Titles That Earn Their Verbs (Required Practice)

Each of the following pairs a view with a proposed title. For each, state whether the title is earned by the view, and if it is not, name which of the three overclaim shapes it exhibits — causal or predictive verb, generalized population, or imported number — and rewrite it as a sentence the view supports.

- A sorted bar chart of average order value by home metro, at order grain, with the chain average drawn. Title: "Suburban customers drive StyleCraft's growth."
- A line chart of the four-week average of weekly contribution by store age for four stores, on a complete store-week scaffold, with a reference line at \$1,769.23. Title: "The suburban format pays back in under a year."
- A 100 percent stacked bar of category revenue share by metro. Title: "Occasionwear is roughly three times the share of suburban revenue that it is of urban."
- A fiscal-year contribution line, actual through November and forecast thereafter, with 80 and 95 percent interval regions on the forecast months only. Title: "Contribution will reach \$1.25 million this year."
- A bar chart of repeat-order share by store type with the chain share drawn. Title: "Suburban repeat-order share is 24 percent against a chain average of 47 percent."

- A scatterplot of items per order against occasionwear share, one mark per store. Title: "Bigger baskets are an occasionwear story."

For the second item, note that all three overclaim shapes are present at once, which is why it is the most instructive title in the list: the verb makes a claim about the future, the noun names a quantity the view does not measure, and the sentence generalizes four observed stores to a format — while also being false of the one that took fifty-seven weeks. State all three, and explain in one sentence why the middle error is the hardest of them to notice.

Then, for the two titles you judged earned, write the alternative title that the same view *also* supports and that would leave a reader with a different impression, and state in one sentence how you would choose between them and what makes that choice a matter for Section 13.15 rather than for Section 13.6.

13.17.4 Subtract, Then Emphasize (Required Practice)

Take any one view you built for Chapter 12's pre-read and convert it into a communication view without changing a single underlying query.

First, write the assertion title. Then list, before touching anything, every element you intend to remove and every element you intend to emphasize, with the reason for each stated in terms of that title. Then make the changes, and record any element you found yourself removing that was not on your list, with a note on why you had not anticipated it.

Finally, produce the before-and-after pair as images and write two sentences: one stating what changed in the reading, in the form Section 12.13.5 established, and one stating what the analytical version could show that the communication version cannot. The second sentence is the point of the exercise. If you cannot name anything that was lost, either the analytical view was overbuilt or the communication view has not been simplified enough, and both are worth discovering.

13.17.5 The Go/No-Go Dashboard and Memo (Homework Submission)

Build the artifact described in Section 13.12, as a Tableau workbook plus a one-page memo. This is the deliverable Project #2 builds on; read the rubric in Appendix F before beginning.

The workbook must contain: the written specification from Lab 13.1 Part A, including the reconciled Table 13.7, the sealed intended takeaway, and the sealed predictions, as a text object or an accompanying page; the store-week scaffold and the certified metrics file, with your own reproduction of every row of Table 13.7 and any mismatch reported as a finding; the seven views of Table 13.8, each rebuilt rather than duplicated, each carrying an assertion title and a grain subtitle; the assembled dashboard at a fixed size, built in containers, with no filter controls and its fixed state declared in words in the source band, one highlight action running from the smaller value set to the larger so that every selection it can produce has a match, one parameter used as navigation with a title that is true in every state it can produce, and alt text on every view; and, on a separate analytical tab, the views and the seven filter controls that came off the canvas.

Four construction requirements are graded specifically because they are the ones most often faked. The predecessor region in View 2 and the interval regions in View 5 must be built from bounds that vary along the horizontal axis, and you must demonstrate it: report the region's upper and lower vertices at three separated store-age weeks and at the first and last forecast months, and show that they differ. View 2's line layer must contain exactly the eight Wave 4 stores and the region must survive the restriction; report the line count and state how you filtered without deleting the band-outline rows. View 5 must carry a horizon-total block whose figures come from the outlook file's horizon rows, and no shading may appear left of the December forecast-to-go seam; report the block's figures beside the sum of the seven monthly bounds and state how far apart they are and why. And every weekly path must be built on the scaffold; report, for one store, `First Observed Coverage Week` and `First 4wk Coverage Week` computed with and without it, and state which of the two moved, which did not, and why the difference is what the scaffold is for.

The memo must contain the four parts of Section 13.10 — recommendation, evidence, uncertainty, reversal conditions — in that order, on one page. Its uncertainty paragraph must state the range still open in store-equivalents in both framings and show the division that produced each, including which denominator you chose and why, must name the \$92,000 footprint as a finance assumption rather than a measurement, and must state the population every figure describes. Its appendix must list every view removed from the canvas with one line apiece on where it went and whether its absence is material to the decision, per C3 of Table 13.6.

Submit both, plus the completed audit: the three transferable points of Table 8.5, all three of Table 12.5 on every view, and all three of Table 13.6 on the artifact, with the C1 sentence count reported — how many sentences you checked and how many you changed.

13.17.6 AI Titles, Prose, and What Disappeared (Homework Submission)

Run the five-step routine of Section 13.11's AI in Practice box, and document it per Appendix D. Submit the following.

First, the title exercise: your request prompt for one view, the twelve candidates with the assistant's verification lines, your written prediction of how many would overclaim and in which shapes, your audit of each candidate against the built view, and your grade of the prediction including both misses and false alarms.

Second, the prose exercise, which is the more revealing of the two. Write your recommendation paragraph yourself, carrying the interval and the reversal conditions. Ask an assistant to make it clearer and more concise for a chief executive. Then diff the two versions clause by clause rather than reading the new one, and report exactly what was removed. State whether each removal is a legitimate compression or a deletion of a finding, and produce a final version that is shorter than yours and contains everything yours contained.

Third, one paragraph on the pattern. Across both exercises, what did the assistant reliably add, and what did it reliably remove? Name the mechanism you think produces each, and state one change to your own prompting that you will make as a result.

Your grade rests on the audit and the diff, not on the assistant's output. An exchange in which the assistant deleted your interval and you caught it is a better submission than one in which it produced a paragraph you accepted.

13.17.7 Where the Line Is (In-Class Discussion)

Each of the following describes a design decision an analyst might make on the go/no-go dashboard. For each, state whether it is emphasis or distortion under the standard of Section 13.15, name the obligation it triggers if any, and state what disclosure would make it acceptable if one exists.

- Placing the coverage view in the largest frame at the upper left, and the repeat-share view small at the bottom.
- Removing the repeat-share view entirely, on the grounds that the contribution already observed reflects it.
- Titling the ramp view "seven of the eight new stores are on track" rather than "one of the eight new stores has tracked below its predecessors since week six."
- Shipping the dashboard with no filter controls at all, and stating the population and the window only in small type in the source band.
- Reporting the coverage milestone as the first single week above the threshold rather than the first four-week average above it, in a memo recommending a ten-year lease.
- Describing the repeat-share movement from 20 to 24 percent as a 20 percent improvement.
- Grouping the eleven smallest product categories into "other" in the mix view, where two of them individually exceed the fourth-largest named category.
- Extending the outlook's horizontal axis six months past the end of the fiscal year so that the interval regions occupy a smaller share of the frame.

For the last item, state what would have to be true about the reader for the choice to matter, and connect your answer to the mechanic Chapter 12 owns that it most resembles.

13.17.8 Ten Minutes with the Chief Executive (In-Class Discussion)

You are the analyst in Section 13.1. It is 9:04 a.m. on December 2. The dashboard is on the screen, you have said your recommendation sentence, and the chief executive has responded: "The merchandising team's numbers last week said suburban revenue jumped last quarter. Yours says the four stores we can observe didn't get their four-week averages up to the weekly fixed-cost threshold until somewhere between weeks 43 and 57. Which is it?"

The numbers she is referring to are the ones from the deck Chapter 12 repaired: a real 4.6 percent quarterly increase, drawn on a truncated axis, built by a junior analyst on the merchandising team. She is not in the room — she was never a pre-read participant. Her manager, the head of merchandising, forwarded the deck with the note "this is great, let's use these," and is sitting at the table.

Prepare a position, and address each of the following explicitly. What do you say in the next thirty seconds, given that each statement is true of what it measures, that the two measure different quantities over different windows, that the discrepancy is partly an encoding artifact rather than a data disagreement, and that the person who forwarded the other chart — and who has argued the composition case for a year — is sitting at the table? Which of Chapter 12's mechanics do you name, and do you name it at all, or do you answer only about your own number? How do you keep the meeting on its decision rather than converting it into a reconciliation exercise, and what do you commit to doing after the meeting instead? What would you do differently if the other deck had come from the chief executive's own prior briefing rather than from a colleague's team? And what would change if the analyst who built it *were* in the room? And finally, a question with no comfortable answer: if you had not built the pre-read, and the only artifact in the room were the other deck, would this decision have gone differently — and what does your answer imply about the value of work whose only visible output is a decision that did not go wrong?

13.18 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.

Assertion title. A chart title that states the finding the view supports, written as a complete sentence with a subject and a verb, in contrast to a label title, which names the fields displayed; because the reader acquires the title before decoding the marks, an assertion title determines what the view is understood to show and must be verifiable from that view's marks in that view's units (adapted from Knafllic, 2015; Kosslyn, 2006).

Audience analysis. The explicit statement, made before an artifact is designed, of who will read it, what decision it must support, how long the reader will give it, what she already knows, and what evidence would change her mind; its output is a set of design constraints rather than a description (adapted from Knafllic, 2015; Few, 2006).

Critique protocol. A structured procedure for feedback on a visual artifact in which the reviewer reports the takeaway, asks the intent, identifies impeding elements, and only then proposes changes, while the author listens without defending; the fixed order is the mechanism, since reporting before learning the intent yields an uncontaminated measurement (adapted from Krug, 2014; Few, 2006).

Dashboard action. An interaction configured on a dashboard by which activity in one view changes another — filtering, highlighting, or navigating to a different sheet — used to link an entity across views or to move a reader from a summary to its supporting evidence on request. A highlight action operates on matching field values and does nothing in a target view that does not contain them; the communicative test is whether the action answers a question the analyst predicted (adapted from Heer & Shneiderman, 2012; Shneiderman, 1996; Tableau, 2026b).

Data-ink and decluttering. Data-ink is the portion of a chart's ink representing data values; non-data-ink is everything else. Decluttering is the systematic removal or de-emphasis of non-

data-ink and duplicated data-ink, undertaken to raise the contrast between what carries meaning and what does not, and retaining the non-data-ink — reference lines, direct labels, annotation, provenance — that helps a reader interpret correctly (adapted from Tufte, 2001; Few, 2012).

Data story. An ordered sequence of views, each carrying a stated point, connected so that the sequence constitutes an argument, typically moving from an accepted situation through a complication to a resolution; its defining property is that the order is chosen by the author and is load-bearing (adapted from Segel & Heer, 2010; Knafllic, 2015).

Exploratory and explanatory artifacts. An exploratory artifact lets a reader investigate by her own path and supports questions the builder did not anticipate; an explanatory artifact delivers a specific finding along a path the builder chose. The distinction is separate from the analysis-communication distinction of Section 12.2, and the most common design error in business visualization is building an exploratory object for a reader who needed an explanatory one (adapted from Knafllic, 2015; Segel & Heer, 2010).

Weekly coverage, observed and four-week-average. First *observed* weekly coverage is the earliest store-age week in which a store's weekly contribution reaches the declared weekly share of its annual fixed footprint; first *four-week-average* coverage is the earliest week in which the trailing four-week average of that contribution reaches it. Both are coverage milestones rather than payback periods, since nothing accumulates, and neither establishes that the store remained above the threshold afterward — a four-week average can be lifted over the line by one strong week and fall back below it. The four-week-average measure is the decision rule in this chapter because it is the less easily gamed of the two, not because it certifies persistence, and the gap between the two is itself informative. Because store age is derived from the opening date, a complete store-week scaffold does not by itself move the observed week; it changes the drawn path and the comparison range, and it is a precondition for the four-week-average measure, whose window cannot be computed on an extract with holes in it (Section 13.12.1; Sections 2.7 and 12.9).

Interval ribbon. A shaded region drawn around a forecast or a comparison set whose upper and lower edges vary along the horizontal axis. Because a Tableau reference band is defined between two values at table, pane, or cell scope and therefore renders at constant height, and because separate continuous measures on Rows occupy separate axes rather than stacking into a region, the ribbon is described upstream as an ordered set of vertices — along the lower bound left to right, back along the upper bound — and drawn with the polygon mark type using the Path shelf. Nested intervals are emitted as edge-to-edge regions rather than layered ones, so that no drawing order can conceal one behind another, and the series line is brought over the regions with **Move Marks to Front** rather than by relying on which axis a pill sits on (Section 13.9; Tableau, 2026a, 2026h, 2026i, 2026j).

KPI hierarchy. The ordering of a deliverable's metrics into tiers by proximity to the decision — primary, supporting, and diagnostic — and the expression of that ordering in the artifact's visual prominence, so that size, position, and emphasis correspond to decision relevance; a

hierarchy stated in the analyst's head but not in the layout is not communicated (adapted from Few, 2006, and the metric and KPI roles of Section 3.5).

Preattentive attributes. Visual properties processed in parallel across much of the visual field before focused attention is directed, so that an object differing from its neighbors in one such property can be detected rapidly and with little increase in search time as the number of neighbors grows. The effect is strongest with a single distinctive feature and homogeneous distractors, and degrades with target-distractor similarity, heterogeneous displays, conjunctions of two attributes, and displays in which several attributes vary at once (adapted from Treisman & Gelade, 1980; Healey & Enns, 2012).

Three-second test. A course heuristic: showing a finished artifact to a representative reader for approximately three seconds, removing it, and asking what it was about. It measures what the artifact communicates before the reader decides to work at it, and a failure is a defect in the artifact rather than in the reader. The formal first-impression procedure in the usability literature is normally run at five seconds (Doncaster, 2014); three is used here as a deliberately demanding executive-dashboard constraint (course heuristic; informed by Doncaster, 2014; Few, 2006; Krug, 2014).

Usability test. A structured observation in which a representative reader who has not seen the artifact is given a task and observed while attempting it, with the analyst recording what she looks at, says, asks, and concludes, and refraining from explaining or guiding; a hesitation or misreading is evidence of a possible mismatch among artifact, task, and audience and is investigated as a candidate defect in the artifact. A reader who has seen an artifact can no longer supply a fresh first impression of it, though she remains available for later task-completion testing (adapted from Krug, 2014; Nielsen & Landauer, 1993).

13.19 Further Readings

Students who want additional background may begin with the following readings. The communication treatments are listed first, the perceptual and experimental evidence second, and the testing practice and tool reference last.

- Knaflitz (2015) for the most direct practitioner's treatment of the material in Sections 13.6 and 13.7 — audience, story, decluttering, and emphasis, worked through with before-and-after examples that are worth studying as pairs rather than reading as advice.
- Few (2006) for dashboard design specifically, and in particular for the argument that a dashboard's defining constraint is the single screen; its catalog of common dashboard failures remains the best diagnostic list available and covers most of what Section 13.1's first draft did wrong.
- Segel and Heer (2010) together with Cairo (2016) for the analytical account of narrative visualization and the broader treatment of visualization as a truth-seeking practice. Read Segel and Heer for the author-driven to reader-driven spectrum that Section 13.3 uses to explain why the go/no-go artifact is a data story wearing a dashboard's clothes; read

Cairo on how much a graphic should show, which is the companion problem to Section 13.15's question of what may be left out.

- Healey and Enns (2012) for the survey of the perception research behind Section 13.7, which explains why emphasis is a scarce budget rather than a technique, where the conjunction penalty comes from, and under what conditions the preattentive effect is weaker than the popular summaries suggest.
- Krug (2014) for usability testing as a practice a non-specialist can actually run, including the discipline of not helping the reader, which is the hardest part of Lab 13.2 and the part the book is most useful on.
- Tableau (2026a–2026k) for the authoritative current build paths behind every instruction in Section 13.12 — reference lines and bands, dashboard actions, filter scope, accessibility and alt text, dashboard sizing and device layouts, dynamic titles, multiple-measure axes, mark properties, mark order, Prep multiple-row calculations, and the forecasting options this chapter deliberately declines to use. Where a printed source and the vendor's documentation disagree about an interface, the documentation governs, and it is the reference an analyst should learn to check first.

13.20 References

Bezos, J. (2018). *2017 letter to shareholders*. Amazon.com, Inc.

<https://www.aboutamazon.com/news/company-news/2017-letter-to-shareholders>

Cairo, A. (2016). *The truthful art: Data, charts, and maps for communication*. New Riders.

Correll, M., & Gleicher, M. (2014). Error bars considered harmful: Exploring alternate encodings for mean and error. *IEEE Transactions on Visualization and Computer Graphics*, 20(12), 2142–2151. <https://doi.org/10.1109/TVCG.2014.2346298>

Doncaster, P. (2014). *The UX five-second rules: Guidelines for user experience design's simplest testing technique*. Morgan Kaufmann.

Few, S. (2006). *Information dashboard design: The effective visual communication of data*. O'Reilly Media.

Few, S. (2012). *Show me the numbers: Designing tables and graphs to enlighten* (2nd ed.). Analytics Press.

Healey, C. G., & Enns, J. T. (2012). Attention and visual memory in visualization and computer graphics. *IEEE Transactions on Visualization and Computer Graphics*, 18(7), 1170–1188. <https://doi.org/10.1109/TVCG.2011.127>

Heer, J., & Shneiderman, B. (2012). Interactive dynamics for visual analysis. *Communications of the ACM*, 55(4), 45–54. <https://doi.org/10.1145/2133806.2133821>

Hullman, J. (2020). Why authors don't visualize uncertainty. *IEEE Transactions on Visualization and Computer Graphics*, 26(1), 130–139. <https://doi.org/10.1109/TVCG.2019.2934287>

- Hullman, J., & Diakopoulos, N. (2011). Visualization rhetoric: Framing effects in narrative visualization. *IEEE Transactions on Visualization and Computer Graphics*, 17(12), 2231–2240. <https://doi.org/10.1109/TVCG.2011.255>
- Knaflic, C. N. (2015). *Storytelling with data: A data visualization guide for business professionals*. Wiley.
- Kosslyn, S. M. (2006). *Graph design for the eye and mind*. Oxford University Press.
- Krug, S. (2014). *Don't make me think, revisited: A common sense approach to web usability* (3rd ed.). New Riders.
- Miller, G. A. (1956). The magical number seven, plus or minus two: Some limits on our capacity for processing information. *Psychological Review*, 63(2), 81–97. <https://doi.org/10.1037/h0043158>
- Nielsen, J., & Landauer, T. K. (1993). A mathematical model of the finding of usability problems. In *Proceedings of the INTERACT '93 and CHI '93 Conference on Human Factors in Computing Systems* (pp. 206–213). Association for Computing Machinery. <https://doi.org/10.1145/169059.169166>
- Segel, E., & Heer, J. (2010). Narrative visualization: Telling stories with data. *IEEE Transactions on Visualization and Computer Graphics*, 16(6), 1139–1148. <https://doi.org/10.1109/TVCG.2010.179>
- Shneiderman, B. (1996). The eyes have it: A task by data type taxonomy for information visualizations. In *Proceedings of the 1996 IEEE Symposium on Visual Languages* (pp. 336–343). IEEE. <https://doi.org/10.1109/VL.1996.545307>
- Spiegelhalter, D., Pearson, M., & Short, I. (2011). Visualizing uncertainty about the future. *Science*, 333(6048), 1393–1400. <https://doi.org/10.1126/science.1191181>
- Tableau. (2026a). *Reference lines, bands, distributions, and boxes*. Tableau Help. Retrieved July 31, 2026, from https://help.tableau.com/current/pro/desktop/en-us/reference_lines.htm
- Tableau. (2026b). *Actions*. Tableau Help. Retrieved July 31, 2026, from <https://help.tableau.com/current/pro/desktop/en-us/actions.htm>
- Tableau. (2026c). *Filter data from your views*. Tableau Help. Retrieved July 31, 2026, from <https://help.tableau.com/current/pro/desktop/en-us/filtering.htm>
- Tableau. (2026d). *Author views for accessibility*. Tableau Help. Retrieved July 31, 2026, from https://help.tableau.com/current/pro/desktop/en-us/accessibility_create_view.htm
- Tableau. (2026e). *Size and lay out your dashboard*. Tableau Help. Retrieved July 31, 2026, from https://help.tableau.com/current/pro/desktop/en-us/dashboards_organize_floatingandtiled.htm

- Tableau. (2026f). *Format individual parts of the view*. Tableau Help. Retrieved July 31, 2026, from https://help.tableau.com/current/pro/desktop/en-us/formatting_specific_titlecaption.htm
- Tableau. (2026g). *Forecast options dialog box*. Tableau Help. Retrieved July 31, 2026, from https://help.tableau.com/current/pro/desktop/en-us/forecast_options.htm
- Tableau. (2026h). *Build a view to compare multiple measures*. Tableau Help. Retrieved July 31, 2026, from https://help.tableau.com/current/pro/desktop/en-us/multiple_measures.htm
- Tableau. (2026i). *Control the appearance of marks in the view*. Tableau Help. Retrieved July 31, 2026, from https://help.tableau.com/current/pro/desktop/en-us/viewparts_marks_markproperties.htm
- Tableau. (2026j). *Move marks to the front or back*. Tableau Help. Retrieved July 31, 2026, from https://help.tableau.com/current/pro/desktop/en-us/move_marks.htm
- Tableau. (2026k). *Create multiple-row calculations in Tableau Prep*. Tableau Help. Retrieved July 31, 2026, from https://help.tableau.com/current/prep/en-us/prep_multirow_calculations.htm
- Treisman, A. M., & Gelade, G. (1980). A feature-integration theory of attention. *Cognitive Psychology*, 12(1), 97–136. [https://doi.org/10.1016/0010-0285\(80\)90005-5](https://doi.org/10.1016/0010-0285(80)90005-5)
- Tufte, E. R. (2001). *The visual display of quantitative information* (2nd ed.). Graphics Press.
- Tversky, A., & Kahneman, D. (1981). The framing of decisions and the psychology of choice. *Science*, 211(4481), 453–458. <https://doi.org/10.1126/science.7455683>