

Appendix E

Data Visualization Checklist

Dr. Jose Mendoza, Academic Director and Clinical Associate Professor

Version 1.0 · August 2026

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

Appendix Information

PURPOSE

This appendix is the last check a chart, an analytical pre-read, an executive dashboard, or a recommendation memo passes before it leaves the analyst's workspace. It consolidates the visual-mechanics obligations of Chapter 12 with the audience, hierarchy, uncertainty, usability, and ethical obligations of Chapter 13 into fourteen checkpoints that can be run rather than remembered, sorts every check into the ones that stop an artifact and the ones that are judgment calls, and ends in a signed record of what failed, what was repaired, and what the analyst decided to accept.

VERSION AND DATE

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

SUGGESTED CITATION

Mendoza, J. (2026). Data visualization checklist. In *Applied business analytics for marketing decision-making: Business analytics and data visualization* (Appendix E, Version 1.0) [Open educational resource]. CC BY 4.0.

LICENSE AND RIGHTS

Except where otherwise noted, this appendix is licensed under a Creative Commons Attribution 4.0 International License. Copyright © 2026 by Jose Mendoza.

Tableau, Tableau Desktop, and Tableau Prep are trademarks of Salesforce, Inc. Google Colab and Google Drive are products 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. "Python" and the Python logos are trademarks or registered trademarks of the Python Software Foundation. Matplotlib and pandas are NumFOCUS fiscally sponsored projects. Product names are used for identification only and do not imply endorsement. StyleCraft Collective is a fictional company created for instruction.

GENERATIVE AI USE

Generative artificial intelligence and other AI-assisted tools were used in the research, writing, revision, and production of this appendix, including outlining, preliminary drafts, revision of prose, 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 appendix's scope, boundaries, and content, and reviewed and approved all AI-assisted material: every check was traced to the chapter section that establishes the obligation it enforces, and every figure quoted from the StyleCraft case was reconciled against the chapter that derives it rather than restated from memory. Responsibility for the accuracy,

originality, and final form of this appendix rests entirely with the author. A fuller statement appears in the front matter of the complete guide.

COMPANION FILES

A printable one-page checklist, an editable Word and Google Docs version, a spreadsheet version carrying the Stop, Caution, and Pass status column, a Tableau-project sign-off form, and an accessible plain-text checklist are in the [Applied Business Analytics companion repository](#).

E.1 What This Checklist Is For

Chapter 12 established that a chart can produce a false impression without containing a false statement, and that the professional standard is therefore not *did I intend to mislead* — a question about the analyst’s interior state that nobody else can check — but *did I check*, which leaves evidence. Chapter 13 added a second question on top of the first: not only whether the artifact shows what the data contains, but whether the reader took away what the artifact actually supports. This appendix exists because both standards fail in exactly the same circumstance, which is a deadline.

An obligation that depends on remembering to care will be honored on quiet weeks and abandoned on the weeks that matter, and the weeks that matter are the ones in which charts reach decision-makers. The merchandising analyst’s deck in Section 12.1 was not built in bad faith; it was built in ninety minutes by a capable person who accepted four software defaults and had no framework for knowing which defaults were dangerous. Making the framework a checklist — read the axis floor, name the grain, name the comparison, run the grayscale check, write the footer — removes it from the domain of virtue and puts it in the domain of process, which is where obligations survive.

CONCEPT

Every Box Checked Is Not a Certification

This checklist can establish that an artifact was inspected against fourteen classes of known failure. It cannot establish that the artifact is correct, that the analysis beneath it is sound, or that the recommendation it supports is the right one. Those are analytical judgments, and a checklist that claimed to make them would be doing the thing this guide has warned against since Chapter 1, which is treating a completed procedure as evidence that thinking occurred.

What a completed checklist is good for is narrower and genuinely valuable. It makes the analyst’s inspection auditable, it makes the accepted defects visible as decisions rather than as oversights, and it converts a recurring argument about whether a bar chart may start at 4,000 into a settled rule that was decided once, in a document, by people who were not under deadline. Treat a run of this checklist as a record of what was checked, and treat the sign-off record of Section E.7 as the artifact that record produces.

Source: Course concept developed for this guide, informed by Sections 12.16 and 13.15.

The division of labor is worth stating once, because several checks ask the analyst to confirm something rather than to compute it, and the difference matters. Table E.2 records what the checklist decides and what it does not.

Table E.2

What the checklist settles and what remains the analyst's

The checklist settles	The analyst decides	Established in
Whether a length-encoded mark begins at zero	Whether the finding survives its honest rendering, and whether there is a finding at all	Sections 12.6, 12.7
Whether the grain of one mark was stated and confirmed	What grain the question requires	Sections 3.3, 12.10
Whether a reference is on the chart and labeled	Which reference the reader needs, and whether it was declared before the results were inspected	Sections 2.7, 12.9
Whether an assertion title is verifiable from the marks beneath it	Which of the defensible sentences to write, and therefore what the reader takes away	Sections 13.6, 13.15
Whether the removals were listed	Whether each removal is a simplification or the deletion of a qualification	Section 13.15
Whether uncertainty is present and in decision units	How much uncertainty the decision can carry, and what would reverse the recommendation	Sections 2.7, 13.9
Whether a reader's takeaway was measured	What to do when the reader's takeaway is not the intended one	Sections 13.2, 13.10
Whether every AI-touched sentence was verified against a view	Whether to keep the sentence	Sections 12.12, 13.11

Read the second column as the reason the first column is short. Almost every entry in it is a decision. Chapters 12 and 13 spend pages establishing, and none of them can be made by inspection. The checklist's job is to make sure the decision was made rather than inherited.

E.2 Two Check Levels, Three Status Labels, and One Sign-Off Record

Two distinct things get labeled in this appendix and confusing them is the most common way a checklist becomes useless. A **level** is a property of the check and never changes: it says whether the obligation is a hard error or a judgment call, and there are two. A **status** is the outcome the analyst records on a particular artifact: it says what happened when the check was run, and there

are three. A single check therefore has one level and, on each artifact, one status. Running the two levels across an artifact produces the third object in this section's title, the sign-off record of Section E.7, which is a document rather than a level.

E.2.1 The Two Check Levels

Level 1 — stop-ship. A Level 1 check identifies either a defect that makes the artifact incorrect, materially misleading, or unusable, or the absence of a mandatory control that makes the audit record invalid. The first kind is recognized by its effect on a reader: she would believe something about the quantities that the data does not support, and would revise that belief on seeing the evidence. The second kind is recognized by what it does to the audit: a decision never written down, a hierarchy never stated, a fault never predicted before the chart was rendered — none of these misleads anyone by itself, and each one makes it impossible to establish afterward that the artifact was checked rather than merely approved. Both kinds are resolved before the artifact circulates, or the artifact does not circulate, and neither is weighed against a deadline.

Level 2 — communication and judgment. A Level 2 check evaluates whether the intended reader can find, understand, and use the evidence. Its failures are real and are graded against a stated reader rather than against a rule, which means a Level 2 check can be failed defensibly: an analyst who decides that a legend is acceptable here, that this window is the right one, or that this view stays on the canvas has made a decision, and recording the decision with its reason is what makes it a decision rather than an omission.

The sign-off record is what running the two levels produces, and Section E.7 specifies it. It carries what was run, what passed, what failed, what was repaired, what was accepted, and who signed. A run of the two levels that leaves no record has inspected an artifact without leaving evidence that it was inspected, which is the state this appendix exists to replace. The record is not a third level, because nothing in it is a check; it is the output the checks are run to produce.

The two levels do not correspond to importance. A Level 2 hierarchy failure on an executive dashboard — the reader cannot tell what to look at first — will more often cost a decision than a Level 1 units label missing from an axis. They correspond to *how the failure is resolved*: a Level 1 failure has a correct answer, and a Level 2 failure has a defensible one.

E.2.2 The Status Labels

Three statuses are recorded, and all three are written as words. Color may accompany them and may never replace them, both because status communicated by color alone fails for a substantial share of readers and because a printed or grayscale copy of a checklist is exactly the copy somebody will consult under deadline (W3C, 2024; Wong, 2011).

Table E.3

The three status labels and what each one obliges

Status	What it means	What it obliges
Pass	The check was run and the artifact satisfies it.	Nothing further. Record it as run; a check recorded as passed without being run is the failure this appendix cannot detect.
Caution	The check was run, the artifact does not satisfy it, and the analyst has decided to accept the defect. Available on Level 2 checks only.	One sentence naming the defect and the reason for accepting it, in the sign-off record. An accepted defect is legitimate; an unrecorded one is not.
Stop	The check was run and the artifact fails it. Mandatory on any failed Level 1 check.	Repair before circulation, a repair note per Section E.6, and a re-run of the checks the repair touched.

Two outcomes sit outside the three statuses and are recorded rather than scored. A checkpoint the routing of Section E.3 does not send to this artifact is marked **not applicable**, with its reason in three words; “not a dashboard” is a reason and a blank cell is not. And a Level 2 check that failed and was then repaired is recorded as **Pass**, with the change listed under communication revisions in the sign-off record, so that the record distinguishes a check that never failed from one that failed and was fixed.

One asymmetry in that table is deliberate and worth naming, because students reach for the wrong label under pressure. **Caution is not available on a Level 1 check.** There is no defensible version of a truncated magnitude bar, an undisclosed filter that changes the conclusion, a per-order claim computed from per-line rows, or a title asserting a number no view contains. If a Level 1 check fails and the artifact ships anyway, the honest record is a Stop that was overridden, with the name of the person who overrode it — which is a different document from a checklist and a conversation worth having before it becomes necessary.

CONCEPT

The Default Is Not a Defense

Four of Chapter 12's five misleading mechanics are available as software defaults or single clicks, and the fifth is what a date filter does when nobody chooses a date. That fact is usually offered in mitigation and is the opposite of one: a defense of the form "the tool chose the axis" concedes precisely the thing that matters, which is that nobody looked.

Several checks in this appendix therefore ask you to confirm that a property was *chosen* rather than that it is *correct* — check 5.4 on axis range and check 9.3 on the default filter state at Level 1, check 3.6 on sort order and check 8.10 on dashboard sizing at Level 2. An inherited default that happens to be right still fails those checks, because the artifact will change and the default will follow it. Axis ranges recompute silently on every filter, measure swap, and table calculation. An axis is not a thing an analyst sets once; it is a thing an analyst reads every time the view changes, and the reading takes two seconds.

Source: Course concept developed for this guide, informed by Sections 12.6, 12.7, and 12.16.

E.3 What Runs on What

Running one hundred and eighty-eight individual checks on every chart would guarantee that none of them is run on anything. What the appendix asks for instead is a rule about *when* each of its two instruments is used, and the rule turns on whether a complete audit of this artifact already exists.

The first release of an artifact, and any graded submission, runs the applicable master checkpoints in full. That run is the artifact's baseline audit, it is what the sign-off record's counts are taken from, and there is no shorter route to it: a defect nobody has looked for is not caught by an instrument that does not ask about it. **Every run after that is a short form.** Once a baseline audit exists, the short form of Section E.5 is the working instrument — for routing attention on a revision, for the routine re-run after a repair, and for the check before recirculating something already audited — and a short-form question that fails, or whose answer is not obvious, sends the analyst back into the master checkpoint it cites.

The division is not a compromise between rigor and time. A short form consolidates several master checks into one question, which is exactly what makes it fast and exactly why it cannot substitute for a first look: a question about axes will catch a truncated bar on a revision, and it will not by itself surface a filter that was never disclosed, because nobody asked. Table E.4 therefore names both — the short form each artifact class uses, and the master checkpoints the baseline audit covers. The timings are for the short form; the baseline audit takes longer and takes it once.

Table E.4

Which instrument runs on which artifact

Artifact	Short form	Baseline audit covers	Conditional additions	Short-form time
Python chart inside a working notebook (Chapters 4–11)	E.5.1	E.1 (1.1, 1.3), E.2, E.3, E.4, E.5, E.6, E.7	E.13 if AI-assisted; E.14 if exported or submitted	Under two minutes
Tableau analytical view or pre-read (Chapter 12)	E.5.2	E.1, E.2, E.3, E.4, E.5, E.6, E.7, E.11 (11.5, 11.11), E.14	E.9 if the workbook carries controls; E.13 if AI-assisted	Ten to fifteen minutes
Executive dashboard or decision artifact (Chapter 13)	E.5.3	E.1 to E.8, E.10, E.11, E.12, E.14	E.9 if interactive; E.13 if AI-assisted	Thirty to forty minutes, plus the reader test
Recommendation memo paired with a visual artifact	E.5.4	E.1 (1.1, 1.2, 1.6), E.2 (2.13, 2.15), E.6, E.10, E.11 (11.1, 11.2), E.12 (12.12)	E.13 if AI-assisted	Ten minutes, plus the third-reader test

Four notes on the routing. The first is that the checkpoints are ordered as they are because the order is the cheapest one: E.2 comes before E.3 because there is no point auditing the chart form of a view built at the wrong grain, and E.12 comes near the end because a reader test on an artifact that has not passed E.1 through E.11 measures the wrong thing. Run them in order.

The second is that two checkpoints are conditional rather than universal, including on a dashboard. Checkpoint E.9 runs only where a reader can change the artifact; a static decision report, a printed board pack, or a PDF export has no interactivity to audit, and marking E.9 not applicable there is correct rather than lax. Checkpoint E.13 runs only where an assistant contributed a specification, a build path, a layout, a title, alternative text, or a sentence. Both are recorded as not applicable with their reason, never omitted silently.

The third is that a pre-read and a dashboard are different objects and the checklist treats them differently, which is the operational form of Section 12.2's distinction. On an analytical artifact the checks protect density and inspectability: a legend is acceptable, a footer on every view is required, and an assertion title is out of place. On an explanatory artifact the checks protect hierarchy and speed of comprehension: the footer consolidates to one source band, direct labels replace the legend, and every title is a sentence carrying an obligation. Where a check differs between the two, it says so.

The fourth is the one students most often get wrong. The checks run on the artifact as it will be *delivered*, not as it exists in the workbook. A view that passes every check inside Tableau

Desktop and fails check 14.2 when exported to PDF has failed, because the PDF is what the reader sees.

E.4 The Master Checklist

Checkpoint E.1 — Decision, Audience, and Artifact Type

Owned by. Sections 12.1 and 12.2 for the two jobs of a visual artifact; Sections 13.2 and 13.3 for the reader and the form; Table 13.1 for the choice among dashboard, report, data story, and memo. The frame point of Table 8.5 transfers here as the question the artifact was built to answer.

Runs on. Any artifact that leaves the analyst’s screen. On a working notebook chart, checks 1.1 and 1.3 alone.

Table E.5

Checkpoint E.1 — decision, audience, and artifact type

ID	Level	The check
1.1	1	The decision the artifact supports is written in one sentence naming the decision-maker, the choice, and the deadline. If the sentence needs an “and,” the artifact is being asked to support two decisions and one of them belongs elsewhere.
1.2	2	The primary reader is named, together with how long she will spend, what she already knows, and what evidence would change her mind.
1.3	1	The artifact’s dominant job is declared — analytical or explanatory — and the design matches the declaration rather than contradicting it.
1.4	2	The form is the right one for that decision and that reader: chart, table, pre-read, dashboard, report, data story, or memo. A one-time decision made in a meeting by a reader new to the material is a data story with a memo attached, whatever was requested.
1.5	1	The intended three-second takeaway is written down and sealed before any test is run. This is the same obligation as check 12.1 and is recorded once.
1.6	2	The action or next question that should follow the reading is stated.
1.7	1	An explanatory artifact was rebuilt rather than inherited. A view carried over from the analytical version keeps every default the analytical version was entitled to and the explanatory version is not.
1.8	2	Where the artifact serves two audiences at once, the dominant job is named rather than averaged, and the checks that follow are run against that job.

ID	Level	The check
1.9	2	An analytical artifact circulated to other readers carries a written reading guide naming the question each view answers, the grain of one mark in each, and the window they share, so that a reader who disagrees knows what she is disagreeing with.

Most common failure. The exploratory view — made in five minutes, full of defaults, correct for its purpose — pasted into a deck because it was already there. Its close relative is the first draft of Section 13.1: eight correct views of equal size in a regular grid, which asserts that the eight matter equally, a claim the analyst does not believe and would not defend if asked.

If it fails. Failures here are structural and cannot be repaired downstream. A well-built object of the wrong kind is not improved by better titles, and an artifact with no declared job will fail Checkpoint E.8 on hierarchy no matter how carefully its individual views are drawn. Return to Section 13.3, choose the form from the decision and the reader, and rebuild.

Checkpoint E.2 — Data, Source, Grain, and Reconciliation

Owned by. Section 3.3 for grain, Section 12.10 for aggregation in the view, Lab 12.1 Part A for reconciliation, and V1 of Table 12.5, which this checkpoint implements. This is the checkpoint on which an artifact is most often quietly wrong.

Runs on. Every artifact, without exception and without abbreviation.

Table E.6

Checkpoint E.2 — data, source, grain, and reconciliation

ID	Level	The check
2.1	1	The source table or certified extract is named on the artifact, not only in the analyst's memory.
2.2	1	The date window is stated, and its start is defensible rather than inherited from a filter nobody set.
2.3	1	Every filter and exclusion, including defaults, is stated in words on the artifact. A control showing a filter's state is a fact about an interface; a caption stating it is a fact about the evidence.
2.4	1	What one mark represents was written down before the view was inspected, and the view's dimensions and the measure's aggregation produce exactly that.
2.5	1	Every rate names its denominator, and the denominator is the eligible base rather than the convenient one.
2.6	1	Any per-order, per-customer, or per-store metric has been verified by hand on one group. Average line value and average order value differ by a factor no software will warn you about.

ID	Level	The check
2.7	1	Totals reconcile to the certified source wherever reconciliation is expected, and any filtered total that no longer matches carries a stated filter.
2.8	1	Distinct counts use the correct key.
2.9	1	A join did not duplicate records: pre- and post-join row counts were compared and the cardinality was enforced.
2.10	1	Absent observations and zero observations are distinguished. A week in which a store was open and sold nothing is a zero; a week before it opened is a null; a line drawn straight across either one is a claim the data does not make.
2.11	1	Missing values were not silently omitted in a way that changes the claim.
2.12	1	A regenerable view has not been replaced by a stale screenshot where current data is required, and the refresh date is visible.
2.13	1	Any reference value is entered as a constant from a certified source rather than computed from the displayed marks. A pooled rate is computed on the pooled denominator and is not the unweighted mean of the displayed rates; a reference computed from the view also moves when the view is filtered, which makes the sentence above it false.
2.14	2	Where a reference legitimately is a descriptive summary of the plotted marks — a median, a mean, a prior period — it is labeled as one, so that a reader knows the line summarizes the marks rather than standing independently of them.
2.15	1	Where an artifact carries figures describing different populations, each figure’s population is named beside it and no figure is used as though it described both.
2.16	1	Every derived date field is reproducible: the window is anchored to the extract rather than to the system clock, and the week-start convention is declared rather than inherited from a data source’s configuration.

Most common failure. The grain slip, and it is the most expensive error in Part III because it is invisible. StyleCraft’s entire suburban finding disappears into a three percent difference when average order value is computed from order-line rows and reappears as a sixty-six percent difference when it is computed at order grain, and there is no error message, no warning, and no visual difference between the correct chart and the incorrect one. Its companion is the mislabeled claim: a chart built on line-grain data is not wrong for being line-grain, it is wrong when the chart or the text around it makes a per-order, per-customer, or per-store claim the grain does not support.

If it fails. Stop. A Level 1 failure in this checkpoint means the artifact answers a different question from the one it says it answers, and no amount of design repairs that. Fix the grain, the denominator, or the reconciliation, then re-run Checkpoints E.3 through E.6, because the numbers in every title have changed.

Checkpoint E.3 — Question, Comparison, and Chart Form

Owned by. Section 12.5 and Table 12.1 for purpose-first selection (Shneiderman, 1996; Munzner, 2014); Section 12.4 for why the choice is consequential; V3 of Table 12.5.

Runs on. Every view, in every artifact.

Table E.7

Checkpoint E.3 — question, comparison, and chart form

ID	Level	The check
3.1	1	The question the view answers is written in one sentence.
3.2	1	The comparison the reader must make is explicit, and it is the comparison the question requires.
3.3	2	The analytic purpose is named — distribution, comparison, trend, relationship, composition, or geographic — before the form was chosen.
3.4	2	The chart form makes that comparison easy, and the wrong-form column of Table 12.1 has been read for this purpose.
3.5	2	A table is used where precise lookup matters more than visual pattern.
3.6	2	Categories are sorted in the order the comparison needs, and the sort was chosen rather than left at the software’s default.
3.7	1	Time runs chronologically, and a line is not drawn across categorical positions, which invents a continuity the data does not have.
3.8	2	Where a time series is drawn as bars, discrete periods are the intended reading. Bars are appropriate for discrete monthly or annual comparisons and discard the continuity the eye reads best when the movement between periods is the finding.
3.9	2	Composition is not shown in a form whose middle categories cannot be ranked, and shares are not compared across segments whose baselines float.
3.10	2	Small multiples were considered wherever an aggregate could hide heterogeneity, and the panel order is reproducible from a field rather than arranged by hand.
3.11	1	The view answers the question it was built for, rather than a different question the software aggregated into.
3.12	1	The units placed on a common scale are comparable. A shared axis is an implicit claim that the things it holds can be compared, and that claim is made by the layout rather than by the data.
3.13	2	Where a purpose required more than one view, the additional views were built rather than the purpose narrowed to fit one.

ID	Level	The check
3.14	1	Where a view carries a claim — an assertion title, a finding in accompanying prose, a recommendation — the question it answers was declared before the view was built. Exploratory views legitimately generate questions rather than answer declared ones; what is not available is producing several views, selecting the one that says something, and presenting it as though its question had been asked in advance, which is the visual form of the multiple-comparisons problem of Section 11.8.
3.15	2	Where a view bins a continuous variable, the bin width was chosen and stated. Binning is a transformation applied before encoding, so it changes the picture without changing the data.

Most common failure. Check 3.12, and it is the one that reverses recommendations. Every aggregate view of StyleCraft’s newest stores supports the sentence “the Wave 4 stores are underperforming,” and the comparison being made — eight stores open four months against nine stores open one to two years — is one no analyst would make in a table and that the chart makes irresistible, because the eye compares whatever is placed side by side on a shared scale and does not ask how old anything is. Change the horizontal axis from calendar time to store age and the recommendation flips, with no filter applied, no row removed, and no measure redefined.

If it fails. Return to Table 12.1 and select from the purpose rather than from the form. Where 3.12 fails, the repair is a different horizontal variable, a facet, or an explicit statement of what the comparison does not control for — and the third option is a legitimate answer only when it is written on the chart.

Checkpoint E.4 — Marks and Encodings

Owned by. Section 12.3 for the grammar, Section 12.4 for the perceptual hierarchy (Cleveland & McGill, 1984; Heer & Bostock, 2010), Section 12.11 for faceting.

Runs on. Every view.

Table E.8

Checkpoint E.4 — marks and encodings

ID	Level	The check
4.1	2	The mark type’s implied reading is the intended one: bars read as discrete comparisons, lines as continuous trends.
4.2	1	The most important comparison is carried by the highest-ranking channel practical — position on a common scale, then position on identical non-aligned scales, then length.
4.3	1	Angle, area, volume, and color intensity are not carrying a magnitude the reader must decode precisely.

ID	Level	The check
4.4	1	Size encodings scale area rather than radius, map zero to zero area, and carry a size legend with labeled reference values.
4.5	2	Shape encodings remain distinguishable at the size the artifact will actually be seen.
4.6	2	Multiple encodings do not contradict one another, and no field is encoded twice without a reason that can be stated.
4.7	1	Three-dimensional effects are absent. There is no analytical case for them.
4.8	2	Legend categories match the field's level of measurement, and direct labels replace the legend wherever the lookup would cost the reader more than the labels cost the canvas.
4.9	2	Faceting's cost was accepted deliberately: every panel shares one scale, and the trade of one rung of the hierarchy for the ability to see many subsets at once was the right one for this question.
4.10	1	Where a magnitude comparison and a spatial question were both wanted, they were split across two views rather than forced into one. A map answers where; it does not answer how much.
4.11	1	A filled or shaded map encodes a rate over a defined area rather than a raw count. A choropleth of counts mostly draws a population map, and a map of a per-store quantity mostly draws where the stores are.

Most common failure. Spending the best channel on the least interesting fact. A chart in which every category is a different saturated hue has used its most attention-grabbing channel to encode something the axis already said, and a chart that puts the decision's magnitude comparison into circle area on a map has spent its worst channel on its most important question.

If it fails. Name the comparison, name the channel carrying it, name that channel's rank, and reassign. The repair is usually a different mark rather than a different palette.

Checkpoint E.5 — Axes, Scales, and Baselines

Owned by. Section 12.6 and Table 12.2 for the obligation by chart form; Section 12.7 and Table 12.3 for the mechanics; Section 12.9 for references (Few, 2012; Tableau, 2026d); the *baseline* point of Table 8.5 and V2 of Table 12.5.

Runs on. Every view. On a working notebook chart this checkpoint and Checkpoint E.2 are the minimum.

Table E.9

Checkpoint E.5 — axes, scales, and baselines

ID	Level	The check
5.1	1	Ordinary magnitude bars begin at zero. A bar encodes value in length and length is read as a proportion, so a truncated bar does not exaggerate, it makes a false statement in the channel readers trust most.
5.2	1	Floating, interval, Gantt, and waterfall bars begin at their explicitly defined reference, and the reference is stated on the chart.
5.3	1	Stacked bars and areas begin at zero, and it is stated whether the segments are values or shares.
5.4	1	Every position-encoded axis range was chosen deliberately and is stated where a reader might be misled. A restricted line-chart axis is honest only if the restriction is visible and defensible; “the software chose it” is not a defense.
5.5	2	A logarithmic scale is used only where the question is multiplicative — a growth rate, a ratio, a proportional difference, or a comparison across orders of magnitude — every plotted value is positive, and the scale type is noted for a reader who will not check it.
5.6	1	Axis units are visible, and the aggregation is named where the number is a sum, an average, or a rate.
5.7	1	Panels meant to be compared carry identical scales, and any dual axis in a comparison view is synchronized.
5.8	1	Dual axes are absent unless the two scales are fixed by definition. Two independent ranges manufacture the coupling the reader will read as a relationship.
5.9	1	The displayed window covers at least one full seasonal cycle in a seasonal business, or the shortfall is disclosed and no seasonal claim is made on the strength of it; and the pre-period is shown wherever an event is being credited with what followed it.
5.10	1	The baseline, target, prior period, control, comparison group, or benchmark the reader needs to interpret the values is on the chart. Where there is nothing to compare against, the title contains no comparative word.
5.11	1	Every reference line is labeled, and an externally supplied decision threshold was declared before the results were inspected. A threshold chosen after the fact is a description, not a benchmark.
5.12	2	A decision threshold, a forecast trigger, or an experimental cut is not presented as a neutral constant.
5.13	1	Every axis was re-read after the last change to the view. Ranges recompute silently when a filter is applied, a measure is swapped, or a table calculation is added.

ID	Level	The check
5.14	1	An index or percent-change view draws its base — 100 or zero percent — as a labeled reference line, and names the base period explicitly.

Most common failure. Axis truncation, and the reason it needs a Level 1 check rather than a reminder is that the distortion survives everything a reader can do about it. Drawn on an axis running from \$240,000 to \$255,000, StyleCraft’s quarterly move from \$241,000 to \$252,000 — a real increase of 4.6 percent — produces a second bar roughly twelve times the height of the first, so that a reader converting the drawing into growth sees something near 1,100 percent. The axis labels are correct throughout. Experimental work finds that the distortion persists when the axis is shown clearly and persists when readers are warned (Pandey et al., 2015; Correll et al., 2020), which is why “the numbers are right there” is not a defense and why the repair has to be structural.

If it fails. Draw length-encoded marks from zero and look at what remains. If the honest rendering shows almost nothing, that is the finding, and the decision about whether there is a finding at all is analysis rather than design. Where a small movement is genuinely the point, move it to a position-encoded view — a line on a restricted, labeled axis, or a percent-change chart drawn from a zero reference — rather than restricting a length.

Checkpoint E.6 — Titles, Labels, Annotation, and Footers

Owned by. Section 12.9 for annotation and the footer discipline; Section 13.6 and Table 13.3 for assertion titles; Section 11.10 for the verbs; C1 of Table 13.6, which this checkpoint implements.

Runs on. Every artifact. On an explanatory artifact this checkpoint is run sentence by sentence and the count is part of the record.

Table E.10

Checkpoint E.6 — titles, labels, annotation, and footers

ID	Level	The check
6.1	2	Analytical views carry descriptive titles appropriate to inquiry; explanatory views carry assertion titles, and only where the assertion is fully supported.
6.2	1	Every assertion title is verifiable from the marks in the view directly beneath it, in that view’s own units, with no other evidence admitted. Cover what you know, look only at the marks, and confirm a reader could do it. Every number in an assertion title is therefore visible in the view or directly derivable from what is visible.
6.3	1	The verb is one the evidence licenses. <i>Rose, is, reached, differs,</i> and <i>is associated with</i> are descriptive verbs a view can support; <i>drives, causes, proves,</i> and <i>will</i> are claims about mechanism or the future that a descriptive view cannot carry however strong the analysis behind it.

ID	Level	The check
6.4	1	Every number in accompanying prose — a subtitle, an annotation, a source-band line, a memo sentence — appears in a view or in the certified metrics file the views are built from, and is traceable to whichever it came from. The stricter rule of check 6.2 governs assertion titles, which a reader verifies against one view and cannot check against a file she does not have.
6.5	1	The population in the sentence is the population in the view. Four stores are four stores and not a format; a statement about physical stores is not a statement about the chain.
6.6	1	Percentage change and percentage-point change are distinguished wherever both are available.
6.7	1	Units appear in the axis, the labels, or the title, and every abbreviation is defined.
6.8	2	Labels are direct wherever they remove a legend lookup, and the labeled set is small enough to remain readable.
6.9	2	Annotation explains a relevant event, break, threshold, or exception rather than decorating the view, and the number of annotations is small enough that they retain their force.
6.10	1	The footer carries source, window, filters, grain, and refresh date whenever the artifact leaves the workbook — under every view on an analytical artifact, once for the screen on a dashboard, with per-view grain moved into tooltips only where nothing load-bearing goes with it.
6.11	1	A metric definition accompanies any number a reader could otherwise misinterpret.
6.12	1	A title over a parameterized or filterable view remains true in every state the control can produce, or the title is built dynamically so that it tracks the control.
6.13	2	The claim-making sentences were counted — every title, subtitle, annotation, source-band line, and line of accompanying prose — and that count is the denominator reported against check 6.2.

Most common failure. The overclaiming title, in three recognizable shapes: a causal or predictive verb substituted for a descriptive one, a population generalized from the observations to the category, and a number imported from somewhere other than the view. All three are more interesting sentences than the true one, which is why they are produced and why they survive review.

If it fails. Reject the sentence rather than repairing it, because repairing a title reintroduces the analyst’s own knowledge into a sentence that has to stand without it. Then notice what the true sentence obliges: the moment a title says “seven of the eight,” the view acquires a duty to show eight things distinguishably, and the title has forced a design change rather than labeling one. That is the technique working in the productive direction.

Checkpoint E.7 — Color and Accessibility

Owned by. Section 12.8 and Table 12.4 for palette type by level of measurement; Section 13.7 for emphasis as a budget and for data-ink discipline (Tufte, 2001; Few, 2012); W3C (2024) for the accessibility requirements the guide adopts as conditions of correctness.

Runs on. Every view. Checks 7.1 through 7.5 are Level 1 in every artifact, including a working notebook chart that will be submitted.

Table E.11

Checkpoint E.7 — color and accessibility

ID	Level	The check
7.1	1	Palette type matches the field’s level of measurement: categorical for unordered groups, sequential for ordered ones and for quantities, diverging only where a meaningful midpoint exists and both distance and direction from it are the point.
7.2	1	Color is not the only visual means by which any essential distinction is conveyed. Position, ordering, direct labels, shape, line pattern, or lightness carries it as well.
7.3	1	The grayscale screening has been run: convert the artifact and confirm the distinctions survive. Collapse means rebuild; survival means one test passed rather than all of them, because grayscale exposes reliance on hue and weak lightness contrast and does not simulate every deficiency.
7.4	1	The artifact has been inspected through a color-vision-deficiency simulator. Roughly one in twelve men of Northern European descent has some form of red–green deficiency, and the deficiency is invisible to everyone else in the room including the person who chose the palette.
7.5	1	Contrast meets a stated criterion at the condition the artifact will actually be read in — a projector in a lit room, a phone, a printed page — rather than on the analyst’s monitor. This guide adopts WCAG 2.2 Success Criterion 1.4.3 for text, at 4.5 to 1 against its background, and 1.4.11 for graphical objects and interface components, at 3 to 1 against adjacent colors (W3C, 2024).
7.6	2	Red-to-green is not used as an unexamined good/bad shorthand, and no unordered rainbow scale carries an ordered quantity. An ordered quantity takes a perceptually ordered palette — sequential where it has no meaningful center, diverging where it does — chosen for the variable and for the comparison the reader must make.
7.7	2	Neutral categories are not assigned unintended positive or negative valence.
7.8	2	Emphasis color is spent rather than spread: one accent for the subject, gray for the context, and an emphasized set small enough that the reader does not have to search inside it.
7.9	2	Background, gridlines, borders, and decoration are de-emphasized or removed, and every retained non-data element is doing a job that can be named. Reference lines, direct labels, annotation, and the footer are non-data ink and are kept.

ID	Level	The check
7.10	2	A brand palette is used for identity and accent, with any sequential ramp derived by varying lightness within one brand hue and a neutral gray as the workhorse. A diverging palette takes two accessible hues around a neutral midpoint, since it cannot be built inside a single hue. Where a brand guideline and an encoding requirement conflict on an analytical artifact, the encoding wins.
7.11	2	A categorical palette carries no more than seven or eight hues. Beyond that, color alone stops being a reliable identifier and the distinction moves to direct labels or to another channel.

Most common failure. An ordered field in a categorical palette, or an unordered field in a sequential ramp, with an accessibility omission alongside it — because nothing in a request for a chart mentions color-vision deficiency and no software default raises it. StyleCraft’s four loyalty tiers in a red-to-green ramp fail three ways at once: they use the two hues most often confused, they carry the order in hue rather than in lightness, and they assert a valence the tiers do not have.

If it fails. Rebuild the palette from the field’s level of measurement rather than adjusting hues, then re-run 7.3 and 7.4. Where color was carrying an essential distinction alone, the general repair is redundant encoding, and direct labels, position, and ordering are better safeguards than any choice of hue.

Checkpoint E.8 — Layout and Visual Hierarchy

Owned by. Sections 13.4 and 13.5 for the KPI hierarchy and its spatial expression (Few, 2006); Section 13.7 for decluttering as the precondition of emphasis (Knaflitz, 2015); Treisman and Gelade (1980) and Healey and Enns (2012) for why emphasis is a budget.

Runs on. Any artifact carrying more than one view — a dashboard, a report page, a pre-read spread.

Table E.12

Checkpoint E.8 — layout and visual hierarchy

ID	Level	The check
8.1	1	The KPI hierarchy was written down before the canvas was opened: one or two primary views, a supporting tier, a diagnostic tier, each defined by the reader’s question it answers.
8.2	1	Prominence matches the hierarchy. The view the decision turns on is the largest object and sits on the reader’s entry path, even when it is the plainer object.
8.3	2	Position and prominence agree rather than fight; a large high-contrast object elsewhere on the canvas does not capture the entry fixation.

ID	Level	The check
8.4	2	Reading order matches the sequence in which the reader must acquire the argument.
8.5	2	Views in the same tier are adjacent and share their formatting; views sharing a scale are aligned, and their axes were confirmed identical rather than assumed.
8.6	1	The number of findings the reader must hold at once was stated before layout and is small. Everything else is available in sequence rather than simultaneously.
8.7	2	Whitespace separates functions rather than filling unused space, and padding was set rather than left at zero.
8.8	1	Every view can be justified by its tier and by the point in the reader's sequence at which she needs it. A view with no justification comes off the canvas and goes into the memo's appendix; it does not become a smaller view.
8.9	1	The case's weakest facts are on the canvas rather than in a backup. A reader who discovers a weakness herself concludes that the analyst either missed it or hid it, and both conclusions are expensive.
8.10	2	The sizing mode matches the consumption condition — fixed at a known projection resolution, device-specific for unknown displays — and the choice was recorded rather than defaulted.
8.11	1	The squint test was run: blur the artifact until the text is unreadable and confirm the intended object still dominates. If it does not, the hierarchy is in the analyst's head and the repair is size and position rather than a bigger title.
8.12	1	The primary message survives at the intended display size, and no essential evidence sits below the fold or behind an unexplained interaction.
8.13	2	The artifact can be read without the analyst narrating every element.
8.14	2	Nothing was solved by shrinking. A canvas of twelve small views is not a decluttered canvas of twelve views; it is an illegible one.
8.15	2	Where the authoring tool offers layout containers — as Tableau does — the canvas is built from nested containers rather than floating objects, so that the tiers hold their relationships when the artifact is resized or something is added late. In a tool without them, the equivalent obligation is that the layout survives a change of size or a late addition without being rebuilt.
8.16	2	The title band carries the decision the screen supports and the date rather than a restatement of the data, and the analyst of record is named on it.

Most common failure. Demoting instead of cutting. The tier structure creates a temptation to solve a crowded canvas by making things smaller, and the result restores the undifferentiated grid at a lower point size. Its companion failure is the emphasized set that has grown: a canvas emphasizing six things has emphasized nothing, because the reader must now search the emphasized set.

If it fails. Declutter against the artifact’s own sentence rather than against a general principle. With the assertion title written down, every element can be asked whether it helps a reader believe that specific sentence, and elements with no defensible answer come off without further argument. Move rather than delete — a view to the memo’s appendix, a filter to the analytical tab, a footnote to a tooltip — so that nothing is lost and the reader’s ten minutes are recovered.

Checkpoint E.9 — Interactivity

Owned by. Section 13.8 and Table 13.5 for what communicates and what conceals; Shneiderman (1996) and Heer and Shneiderman (2012) for the interaction types; Tableau (2026a, 2026c) for the mechanics of actions and filter scope.

Runs on. Any artifact a reader can change: a dashboard, a published workbook, an interactive report.

Table E.13

Checkpoint E.9 — interactivity

ID	Level	The check
9.1	1	Every finding the artifact must communicate is visible in the default state. Interactivity may add depth to a finding that has already landed; it may not deliver one.
9.2	1	Each control exists to answer a reader question the analyst can name.
9.3	1	The state the reader meets is disclosed in words on the canvas rather than shown only in a control. Where a screen that could carry controls carries none, it says so, because a screen with nothing to adjust otherwise makes no claim about its own completeness.
9.4	1	No control can falsify a sentence on the same screen. Where no safe scope exists, the control comes off and the state is declared in words instead.
9.5	1	Filter scope was chosen deliberately and stated. A filter applied to every worksheet on a data source will move views whose titles quote chain-wide figures and whose references were computed on another population.
9.6	1	A filter written against a field that only one kind of row possesses does not silently delete the other kind.
9.7	1	No filter or action combination produces an empty, partial, or nonsensical state without explanation — a view emptied of its marks, a frame whose count no longer matches the sentence above it, or a region deleted while its caption remains.
9.8	2	The reader can return to the default state, and the route back is discoverable.
9.9	1	Cross-filtering does not silently change a denominator.

ID	Level	The check
9.10	2	Where one view's selection drives another — a Tableau highlight action, a cross-filter, a linked selection — the interaction runs from the smaller value set to the larger, every target contains the values the source passes it, and each was tested and observed to fire. An interaction that does nothing reads as a broken artifact rather than as an empty match.
9.11	1	Nothing load-bearing lives in a tooltip. A qualification that makes a headline defensible belongs on the canvas.
9.12	2	Navigation actions carry a visible cue that the deeper view exists.
9.13	1	The static export still contains the essential claim and its context.
9.14	2	The interaction inventory is complete: for every interaction, the reader question it answers, what a reader who never triggers it misses, and confirmation that it fires. Where the answer to the second names a finding, check 9.1 has failed and the finding moves onto the canvas.

Most common failure. The presence of many controls, which is itself a message, and the message is “this is a tool for you to use” — the wrong message for a ten-minute reader making a one-time decision. Underneath it sits the assumption that defeats most dashboards: the analyst who built the artifact will click everything, a colleague will click a little, and an executive in a meeting will click nothing, so the analyst reliably overestimates how much of her design the audience will ever see.

If it fails. Apply the first row of Table 13.5 honestly to the specific artifact. A filter communicates when it narrows a population the reader already understands *and* cannot falsify anything on the screen; where the second condition cannot be met at any scope, the honest number of controls is zero, and the state moves into the source band. Removing a control is a communication decision rather than a simplification, and it belongs in the specification with its reasoning.

Checkpoint E.10 — Uncertainty and Decision Conditions

Owned by. Section 13.9 for the executive treatment; Chapters 10 and 11 for the intervals themselves; Section 2.7 for asymmetric cost; the *error in decision units* point of Table 8.5 applied to a range rather than to a point. Correll and Gleicher (2014) and Spiegelhalter et al. (2011) supply the encoding evidence, and Tableau (2026d) is why check 10.6 cannot be satisfied with a reference band.

Runs on. Any artifact carrying a forecast, an estimated effect, a threshold, or a recommendation.

Table E.14

Checkpoint E.10 — uncertainty and decision conditions

ID	Level	The check
10.1	1	Material uncertainty is present wherever the decision depends on it. A point estimate presented alone transfers risk from the analyst's credibility to the organization's balance sheet.
10.2	1	The interval is named correctly — prediction interval, confidence interval, simulated range — and its level is stated.
10.3	1	The interval is expressed in decision units, and the conversion is shown on the canvas in steps with its denominator named rather than asserted as a result.
10.4	1	Where a conversion would require an assumption that cannot be established, what is measured or forecast was changed so the assumption is not needed; where that is impossible, the assumption is disclosed where a reader will find it rather than buried in an arithmetic step.
10.5	1	Where both framings are used, both are given: the end-to-end width of the range and the spread around the central estimate. Quoting the first alone overstates how uncertain the estimate is.
10.6	1	Interval bounds exist as values at the display's own grain, computed upstream rather than requested from the drawing tool. A region whose edges move across the horizontal axis is a shape built from those values — in Tableau a polygon over computed vertices, since a reference band renders at constant height; in a plotting library, a fill between two series.
10.7	1	An interval over an aggregate is not produced by summing the bounds of its component periods, which overstates the width. Where the interval comes from a simulation, the aggregate is formed on each path and the quantiles are taken from the aggregated totals; where it comes from another method, the aggregate interval is derived by that method's own rule and the rule is named.
10.8	1	Periods that have already resolved are shown as actuals outside the interval, with the seam marked and the forecast origin stated. Presenting elapsed months inside a band prices risk that no longer exists.
10.9	2	False precision has been removed and every assumption is named.
10.10	1	The range is not visually compressed — by an over-wide axis, a long empty tail, or a scale chosen after the interval was known — so as to appear more certain than it is.
10.11	1	The artifact does not imply that an observational relationship is an experimental effect.
10.12	1	The complete decision artifact carries all four parts: the recommendation, the evidence, the uncertainty, and the conditions under which the recommendation would change.

ID	Level	The check
10.13	2	Where the cost of being wrong is asymmetric, the asymmetry is stated in comparable units, with both subtractions shown rather than two figures set against each other that are not comparable.
10.14	1	Every assumption supplied rather than measured — a planning figure, a cost, a capacity — is attributed to its source on the artifact rather than presented as something the data produced.
10.15	1	A recommendation states the rule before any count, and any count it carries is reconcilable to a table the reader can be shown. A count that cannot be reproduced from a stated rule is an opinion with a number attached.
10.16	2	The reversal conditions are a small, enumerable set — two or three is the working shape — each with a date or a threshold attached, and at least one names what would be re-estimated when the underlying data next refreshes.
10.17	2	Where part of a forecast horizon has already resolved, the model’s prediction for those periods was compared against what was reported, and the comparison is disclosed as evidence for or against the range on the periods still open.

Most common failure. The interval that disappears during editing, and it is worth predicting explicitly because it is so reliable: 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. The human version is the same move made by the analyst, on the grounds that a range will be read as incompetence.

If it fails. The alternative to a hedge is not confidence; it is the 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 what converts an interval into a monitoring plan, and it is what makes a decision made under uncertainty cheap to unmake.

Checkpoint E.11 — Persuasion, Omission, and Ethical Review

Owned by. Section 12.16 for the disclosure standard; Section 13.15 for the boundary between emphasis and distortion; C3 of Table 13.6, which this checkpoint implements; Hullman and Diakopoulos (2011) for the framing effects, Tversky and Kahneman (1981) for the framing result.

Runs on. Any artifact built for an audience — which is every artifact that is not a private working chart.

Table E.15

Checkpoint E.11 — persuasion, omission, and ethical review

ID	Level	The check
11.1	1	The counterfactual has been asked and answered in writing: would this reader revise her conclusion if she saw everything the analyst saw? A yes means the compression or the emphasis requires repair or disclosure.
11.2	1	Every removal made for the audience is listed — views cut, categories grouped, windows shortened, filters defaulted or scoped, qualifications moved into tooltips — and each is sorted into a legitimate simplification or the deletion of a qualification.
11.3	1	Grouping into “other” does not hide a category whose individual value would change a ranking or a conclusion.
11.4	1	The quantity is presented in a form that does not select the reader’s reaction, and an equally true alternative framing was considered rather than avoided.
11.5	1	The views shown represent the evidence rather than the favorable part of it, and inconvenient findings remain available or are disclosed.
11.6	1	A title does not select the favorable reading where the same view equally supports an unfavorable one that the decision turns on.
11.7	1	Axis, canvas, and scale choices do not minimize uncertainty or adverse outcomes.
11.8	2	The artifact would survive its own audience being fully informed: hand the reader the full evidence and an afternoon, and expect the same decision. An artifact that depends on the reader not having that afternoon is an argument requiring an uninformed audience.
11.9	2	The standard applied was <i>did I check</i> , which leaves evidence, rather than <i>did I intend</i> , which cannot be audited by anyone.
11.10	2	Emphasis was not abandoned in the name of neutrality. A screen with no hierarchy is not a neutral screen; it is one whose hierarchy was set by the order the views were built in, which still produces a conclusion in the reader’s head and is unaccountable for it.
11.11	2	The views that were built and rejected are listed, with one line apiece on why, so that a reader can see the analyst’s search rather than only its result.
11.12	1	A case singled out for attention was selected by a rule declared in advance rather than chosen after inspection and then explained, which is post hoc hunting in a different costume.

Most common failure. The same sentence — “it was cluttering the screen” — is available to describe a view removed because the memo carries it better and a view removed because it complicates the recommendation. That is precisely why check 11.2 asks for the sorting to be done deliberately and in writing rather than left to the analyst’s sense of her own motives.

If it fails. Anything that survives the counterfactual with a yes goes back on the artifact, at whatever size and however inconvenient. Where the material genuinely will not fit, it goes into the memo's appendix and the artifact says on its face that it is there.

CONCEPT

Emphasis Directs Attention; Distortion Directs Belief

Section 13.15 draws the distinction this checkpoint enforces and is where the argument belongs; what the checklist adds is the order in which to run it. Ask check 11.1 first, because it is the only test that reaches a defect assembled entirely out of accurate elements. Then run 11.3, 11.4, and 11.5, which are the three moves — aggregation, framing, and selection — through which that assembly is actually performed.

The reason the order matters is that 11.3 through 11.5 can each be answered yes in good conscience while 11.1 is answered no, and the reverse never happens.

Source: Course concept developed for this guide, informed by Section 13.15, Hullman and Diakopoulos (2011), and Tversky and Kahneman (1981).

Checkpoint E.12 — Usability Testing

Owned by. Section 13.2 for the three-second test as a course heuristic; Section 13.10 for the protocol and the critique exchange; C2 of Table 13.6. Krug (2014) and Doncaster (2014) supply the technique; Nielsen and Landauer (1993) supply the caution about how much a small test establishes.

Runs on. Explanatory dashboards, decision artifacts, and major project deliverables. Not on a working notebook chart, and not on a pre-read whose readers are the people who will interrogate it.

Table E.16

Checkpoint E.12 — usability testing

ID	Level	The check
12.1	1	The intended takeaway was written down and sealed before the test. A takeaway written afterward is not a prediction and the protocol has no value without one.
12.2	2	Three further predictions were sealed: what the reader's eye lands on first, the first question she asks, and the element most likely to be misread.
12.3	1	The reader is representative and has not seen the artifact. A colleague who built the same thing last week knows where to look and is a test of nothing.
12.4	1	The artifact was shown for approximately three seconds, removed, and the reader's takeaway recorded verbatim rather than paraphrased.

ID	Level	The check
12.5	1	A ninety-second task followed — “tell me what you would decide and why” — with reading order, questions, hesitations, and the stated decision recorded.
12.6	1	The analyst said nothing during the task: no explanation, no guidance, no answering of questions until the time was up.
12.7	2	A mismatch in the <i>subject</i> of the reader’s sentence was investigated as a hierarchy failure first, since size and position are the cheapest repair and the most common cause.
12.8	2	A mismatch in the <i>verb</i> was investigated as a title or emphasis failure first, and repaired in language and color before anything structural was changed.
12.9	2	Views the reader skipped entirely were distinguished from views that failed to persuade. A skipped view was never read, and the two have different repairs.
12.10	1	Material mismatches produced a revision and a retest on a fresh reader; accepted defects were recorded with one sentence of reason each.
12.11	1	The result is reported as a formative check on one or two readers, which finds gross hierarchy and language failures, rather than as validation that the artifact works for its actual audience.
12.12	2	The memo alone, without the artifact, was handed to a third reader who could state the recommendation and the basis for it.
12.13	2	Where a peer critique was run, it followed the protocol in order — report the apparent takeaway, ask what decision the artifact supports, identify the specific elements that impede reading, and only then propose changes — with the author listening without defending until the sequence is complete.
12.14	2	Every piece of feedback declined is recorded with one sentence of reason, alongside the changes that were made and the evidence that prompted each.

Most common failure. Talking. The reflex to explain a view the moment the reader misreads it destroys the measurement and is nearly irresistible, which is why check 12.1 exists: writing the prediction down first is what puts something at stake in staying quiet.

If it fails. Revise the artifact, not the reader. The reader’s response is the measurement rather than the error, and a reader who says the artifact is about the new stores doing badly has told you something true about your hierarchy that no amount of inspection would have surfaced.

Checkpoint E.13 — AI-Drafted Visual or Communication Audit

Owned by. Section 12.12 for the chart-drafting failure modes and the five-step routine; Section 13.11 for the language failure modes; Appendix C, Templates C.16 through C.18, for the prompts; Appendix D for the record.

Runs on. Any artifact whose specification, build path, layout, title, alt text, or accompanying prose an assistant contributed to. Run it in addition to the other checkpoints, never instead of them.

Table E.17

Checkpoint E.13 — AI-drafted visual or communication audit

ID	Level	The check
13.1	1	The faults were predicted from the specification, in writing, before the view was rendered. A rendered chart is persuasive in a way its specification is not, and the human tendency is to evaluate whether it is attractive rather than whether it is right.
13.2	1	The specification was built exactly as written, faults included, so that the faults can be audited rather than corrected invisibly. This applies in an instructional or otherwise safe environment; a specification that would create a privacy, security, cost, or irreversible operational exposure is audited on paper and never executed.
13.3	1	No field, filter, value, or business fact was invented. Every column referenced exists in the schema that was supplied.
13.4	1	The grain and the aggregation were stated by the analyst rather than chosen by the assistant.
13.5	1	Software and assistant defaults were inspected individually: axis range, sort order, palette, aggregation, filter state, mark type.
13.6	1	Candidate titles were verified against the built view one sentence at a time and rejected on first failure rather than repaired.
13.7	1	No causal or predictive verb was introduced at the narration step, where assistants reintroduce it even when the numbers are clean.
13.8	1	No interval, threshold, qualification, or reversal condition was removed during a rewrite. The two versions were diffed clause by clause rather than the new one read.
13.9	1	A proposed layout implements the analyst’s KPI hierarchy, supplied as an input, rather than a conventional grid the assistant has no basis for.
13.10	1	No connective sentence asserts a relationship between two views that no view establishes. It reads as synthesis; it is confabulation, and only the traceability rule catches it.
13.11	2	Alt text, footer content, contrast, and filter-state disclosure were requested explicitly, because nothing in a request for a headline or a layout will produce any of them unasked.
13.12	1	The prediction was graded: both the faults missed and the faults predicted that did not occur are recorded, since a systematic error in the analyst’s model of the assistant is worth more than any single caught title.
13.13	1	The exchange is documented per Appendix D, including the output that was rejected.

ID	Level	The check
13.14	1	The analyst can explain and defend every encoding and every sentence in the artifact. An assistant can write a sentence; only the analyst can sign it.
13.15	1	Every candidate sentence arrived with its own verification line naming the marks a reader would have to look at to confirm it. A candidate whose verification line is vague is a candidate that cannot be verified.
13.16	2	The surviving title was shown on its view to someone who had seen neither, and her paraphrase was the claim the analyst intended rather than the stronger one next door.

Most common failure. Two, one per medium. On a chart, the axis truncated on request: an assistant asked to “make the growth clear” will restrict the axis without comment, because the instruction was followed. On a sentence, the overclaim, because a more interesting sentence is better writing and the assistant is optimizing for writing.

If it fails. Repair one mechanic at a time against the checkpoint that owns it, writing a repair note per Section E.6 for each. Where the failure is a sentence, discard it and take a different candidate rather than editing the one that overclaimed.

Checkpoint E.14 — Export and Delivery

Owned by. Section 12.9 for the footer discipline, Section 13.5 for sizing, and Lab 13.1 Part C for alt text; Tableau (2026b, 2026e) for the mechanics. The delivery checks themselves are this appendix’s own: no chapter section owns the export, because in both chapters the artifact is inspected where it is built.

Runs on. Every artifact at the moment it is delivered, in the form it is delivered in.

Table E.18

Checkpoint E.14 — export and delivery

ID	Level	The check
14.1	1	Final dimensions match the delivery channel, and text is readable at actual size in the room, on the page, or on the device where it will be read.
14.2	1	The static export preserves labels, legends, reference lines, annotations, and the footer.
14.3	1	Interactive links open with the intended permissions and in the intended default state.
14.4	2	The workbook, notebook, or query accompanies the artifact wherever reproducibility is required.
14.5	1	The file name and version are unambiguous, and the version on the artifact matches the version in the sign-off record.

ID	Level	The check
14.6	1	The data refresh date is visible wherever recency matters.
14.7	1	Alternative text describes the view’s purpose, population, units, main pattern, and any material exception — not the field names, which is what the software supplies by default.
14.8	1	Every recipient is authorized for everything the delivery carries. Labels, tooltips, annotations, file names, and the extract behind the artifact expose no protected field and no more detail than the decision requires, permissions on any interactive link match the intended audience, and nothing travels that the analyst is not authorized to disclose to the people who will receive it.
14.9	2	A decision-supporting artifact travels with its written recommendation rather than alone.
14.10	2	The sign-off record of Section E.7 is complete, and the analyst of record is named on the artifact or travels with it.

Most common failure. The artifact that passed every check inside the authoring tool and lost half of them on the way out: a legend clipped in the PDF, a tooltip carrying the grain statement that the printed copy cannot show, alt text still naming fields, an interactive link that opens for the analyst and not for the recipient.

If it fails. Repair and re-export, then check the export rather than the source. The only version that matters is the one the reader receives.

E.5 Artifact Short Forms

These four forms are the working instrument for every run after the baseline audit, on the rule stated in Section E.3: the first release of an artifact and any graded submission run the applicable master checkpoints in full, and every run after that is a short form. Each form is a front end to the master checklist rather than a substitute for it. Every item consolidates one or more master checks and cites them by identifier, so the two can never disagree: an item is satisfied when its cited checks are satisfied. The forms carry no level column, because the level belongs to the master check — an item citing 5.1 and 5.4 inherits Level 1 from both, and an item citing 3.10 inherits Level 2.

A short form is deliberately shorter than the checkpoints behind it. Its questions consolidate, and consolidation loses resolution: a form asks about axes rather than about each of the fourteen axis checks, and a defect that no cited question would surface is a defect the form will not find. That is the cost of speed and the reason the baseline audit is not optional. On a revision it is the right trade, because the baseline audit has already looked everywhere once and the short form is looking at what changed.

When a short-form item is satisfied on inspection, record it and move on. When it fails, or when its answer is not obvious, open the master checkpoint it cites, work the individual checks, and record the outcome against those identifiers rather than against the short-form question —

because a repair note, a re-run, and a chapter cross-reference all need the individual check, and “the axis question failed” names none of them.

E.5.1 Single Chart Sign-Off

For a chart in a working notebook, a figure in a report, or one view built in Tableau. Under two minutes.

Table E.19

Short form E.5.1 — single chart

Ask	From
What decision or question does this view support, and is its job analytical or explanatory?	1.1, 1.3
What question does this view answer, and what comparison must the reader make?	3.1, 3.2
What does one mark represent, and do the dimensions and aggregation produce that?	2.4
Is the channel carrying the main comparison the best one available?	4.2
Does every length-encoded mark begin at zero, and was every axis range chosen and stated?	5.1, 5.4
Is the reference the reader needs on the chart, is it labeled, and was it entered from a certified source?	5.10, 5.11, 2.13
Does the palette match the field, and does the chart survive grayscale?	7.1, 7.3
Does the title claim only what the marks support, with its units and its verb?	6.2, 6.3, 6.7
Does the footer carry source, window, filters, grain, and refresh date?	6.10
Was the sort chosen rather than inherited?	3.6

E.5.2 Analytical Pre-Read Sign-Off

For a set of views built to be argued with. Its governing property is inspectability: a reader who disagrees must be able to see exactly what she is disagreeing with.

Table E.20

Short form E.5.2 — analytical pre-read

Ask	From
Is the artifact's job declared as analytical, and does its density serve inspection rather than speed?	1.3
Is there a written reading guide naming the question each view answers, its grain, and the shared window?	1.9
Does every view state its own grain, window, filters, and source?	2.1, 2.2, 2.3, 2.4, 6.10
Do the totals reconcile, and has one per-unit metric been hand-checked?	2.6, 2.7
Are the units placed on any shared scale genuinely comparable?	3.12
Are the pre-periods, comparison groups, and full seasonal cycles present?	5.9, 5.10
Were small multiples considered wherever an aggregate could hide heterogeneity?	3.10
Are definitions and scales consistent across views, and every reference a certified constant?	2.13, 5.7, 6.11
Are the titles descriptive rather than assertive, so that the reader forms her own conclusion?	6.1
Is the inconvenient view present, and are the charts that were built and rejected listed?	11.5, 11.11
Does the delivered export carry the labels, references, and footers the workbook shows?	14.2, 14.6

E.5.3 Executive Dashboard Sign-Off

For a single-screen decision artifact. This form is the compressed version of the full fourteen checkpoints and does not replace them on a first build; it is what you run on the second and third revisions.

Table E.21

Short form E.5.3 — executive dashboard

Ask	From
Is the decision written in one sentence, with the reader, the choice, and the deadline?	1.1, 1.2

Ask	From
Was every view rebuilt rather than inherited from the analytical version?	1.7
Is the KPI hierarchy written down, and does prominence match it under the squint test?	8.1, 8.2, 8.11
Does the title band carry the decision and the analyst's name?	8.16
Is every sentence on the screen verifiable from the marks beneath it, in its own units, with a licensed verb and no imported number?	6.2, 6.3, 6.4, 6.5
Are the case's weakest facts on the canvas?	8.9
Is the uncertainty present, in decision units, with the conversion shown and its denominator named?	10.1, 10.3
Are the assumptions attributed to their sources, and the reversal conditions stated with dates or thresholds?	10.14, 10.16
Is the default state disclosed in words, and can no control falsify a sentence on the same screen?	9.3, 9.4
Is every finding visible without interaction, and does no control produce an unexplained empty state?	9.1, 9.7
Was the three-second test run on a fresh reader, and did her takeaway match the sealed one?	12.1, 12.4
Is the removal list written and sorted?	11.2
Does the export carry labels, alt text, refresh date, and no restricted data?	14.2, 14.7, 14.8

E.5.4 Recommendation Memo Sign-Off

For the one-page written recommendation that accompanies a decision artifact. A memo is checked against the same sentence-level standard as a title, because it makes the same kind of claim in a more durable form.

Table E.22

Short form E.5.4 — recommendation memo

Ask	From
Does it open with a recommendation — an action or a rule — rather than a finding wearing a recommendation's clothes?	10.12

Ask	From
Is the recommendation a rule before it is a count, and can any count be reconciled to a table in the appendix?	10.15
Does every evidence sentence name the view it comes from and state its finding in decision units?	6.2, 10.3
Is every verb one the evidence licenses, and every population the one the evidence describes?	6.3, 6.5
Do the memo's numbers agree with the artifact's, sentence by sentence?	6.4
Is the uncertainty in the reader's units, with the arithmetic and its denominator shown?	10.3, 10.5
Are the scope limits stated — the populations the evidence does not cover, and the false precision removed?	2.15, 10.9
Is every assumption supplied rather than measured attributed to its source?	10.14
Are there two or three reversal conditions, each with a date or a threshold attached?	10.16
Are the removal list and the Appendix D documentation in the appendix?	11.2, 13.13
Handed to a third reader without the artifact, can she state the recommendation and its basis?	12.12

E.6 The Repair Note

Every Level 1 failure produces a repair, and every repair produces one sentence. The sentence records what changed in the *reading*, not what changed in the chart, and the distinction is the whole point of the instrument: “I set the axis to include zero” is a description of an action, while “a reader now sees a quarter that moved slightly rather than a quarter that appeared to jump twelvefold” is a description of a consequence, and only the second tells anyone whether the repair mattered.

Three uses follow from writing them that way. The note is what goes into the memo you send a colleague whose chart you corrected, where naming the mechanic rather than the person is what makes the message survivable. It is the evidence that a repair was worth making, which matters when a stakeholder asks why an artifact was delayed by a day. And it is the record that lets a later reader distinguish a repaired artifact from one that never needed repair, which is the difference between a checklist that was run and a checklist that was filed.

Table E.23

Repair notes, in the form they should be written

Check	The repair	What changed in the reading
5.1	Redrew the quarterly bars from a zero baseline.	A reader now sees a quarter that moved slightly, rather than a quarter that appeared to jump twelvefold.
5.8	Replaced the dual axis with two panels sharing the time axis, and added the urban series as the comparison group.	A reader now sees two series she must compare deliberately, rather than one apparent relationship she was invited to accept.
3.9	Rebuilt the seven-category pie as a sorted bar, then as a 100 percent stacked bar split by metro.	A reader can now rank the middle categories, which the pie made impossible, and can see the share comparison the pie was pretending to answer.
7.1	Replaced the red-to-green tier palette with a single-hue sequential ramp ordered None to Backstage.	A reader now sees four ordered tiers, rather than two tiers she cannot distinguish and a good/bad judgment the tiers do not carry.
2.4	Repointed the sheet at the order-grain extract and relabeled the measure.	A reader now sees average order value, rather than average line value under a label that promised the first.
6.2	Replaced an assertion title with the sentence the view's marks support.	A reader now takes away a claim she can verify from the chart in front of her, rather than one she must take on trust.

The last column of that table is the deliverable in a way the middle column is not, and it is also the column students find hardest, because writing it requires deciding what the reader believed before and what she believes now. An analyst who cannot state the difference has not established that the repair was necessary.

E.7 The Sign-Off Record

The record is one page and it is the appendix's output. It carries the identity of what was reviewed, evidence that every applicable check was run, the failures found and their disposition, the judgment calls that were accepted, the usability evidence where the artifact required it, and a name. Where the artifact will be submitted for a grade, the record is submitted with it; where it will be circulated at work, the record stays with the workbook.

One design decision in the record is worth stating before the fields. A one-page sign-off cannot carry an outcome line for every applicable check, and it does not try to; the completed short form, or the companion spreadsheet where the master checklist was worked, is where the individual outcomes live. What the record carries instead is the **count line** — how many distinct checks were applicable, how many passed, how many were accepted as Cautions, how many Stops were repaired, and how many remain unresolved — together with the identifier of the completed checklist the counts came from. The counts are in unique check identifiers rather than

in check executions, because an artifact of four views runs the same check four times and a one-page record that counted executions would tell a reader nothing she could reconcile. Without that link a sign-off proves that problems were recorded and does not prove that the checks were run, which is a materially weaker claim and the one most sign-off documents in practice actually make.

Table E.24

The sign-off record

Field	What it holds
Artifact title	The name a reader will see, not the file name.
File and version	The delivered file, its version, and the format it was delivered in.
Data refresh date	The date of the extract or the certified source behind the artifact.
Population and window	The population every store-level or unit-level figure describes, and the window, in the words used on the artifact.
Instrument run	Whether this was the baseline audit — the applicable master checkpoints run in full — or a short-form re-run against an existing baseline, and which of the four forms of Section E.5 was used as the front end.
Completed checklist	The file name, spreadsheet tab, or record identifier of the completed checklist the counts below are taken from. A sign-off without this line asserts its counts rather than evidencing them.
Checkpoints run	The checkpoints run, and the checkpoints recorded as not applicable with a three-word reason each.
Unique applicable check IDs	How many distinct checks applied to this artifact, given the routing of Section E.3 and the checkpoints marked not applicable. Counted once for the artifact, not once per view: the individual instances live on the completed checklist.
Unique check IDs passed	How many distinct applicable checks were run and satisfied on every instance.
Unique Level 2 check IDs accepted	How many distinct Level 2 checks were accepted as Cautions. Must equal the number of entries in the accepted-defects field.
Unique Level 1 check IDs repaired	How many distinct Level 1 checks failed anywhere in the artifact and were repaired everywhere they failed. Must equal the number of entries in the stop-ship field. One check failing on four views is one identifier, four instances, and potentially four repair notes.
Unique Level 1 check IDs unresolved	How many distinct Level 1 checks failed and were not repaired. Any number above zero means the artifact has not been released; if it was released anyway, the override and the name of the person who authorized it go here.

Field	What it holds
Stop-ship failures found	Every Level 1 failure by identifier, with its repair note per Section E.6.
Communication revisions	Every Level 2 change made, by identifier.
Accepted defects	Every Level 2 check recorded as Caution, by identifier, with one sentence of reason.
Sentence count	For an explanatory artifact: how many claim-making sentences were checked against 6.2, and how many were changed.
Removal list	The C3 disclosure of check 11.2: what came off, and whether each removal is material to the decision.
Usability evidence	The sealed intended takeaway, the reader's verbatim three-second takeaway, the reading order observed, and the revisions that followed.
AI-use record	The Appendix D record identifier for every exchange that touched the artifact, and what was rejected.
Remaining limitations	What the artifact still does not establish, in the analyst's own words.
Analyst of record	Name and date. The person who circulates the artifact owns every encoding in it.
Reviewer	Name and date, where a second reader ran the checklist.

Three fields in that table are the ones usually left blank and are the ones worth insisting on. **Completed checklist** is what converts the four counts from a claim into a reference; the counts and the completed form must reconcile, and reconciling them takes ten seconds and is the only check anyone can run on the record itself. **Accepted defects** is what distinguishes a decision from an oversight, and its value is entirely in the sentence: a legend retained because the audience is technical and the canvas is crowded is a judgment; a legend retained because nobody looked is not, and the record is the only place the difference is visible. **Remaining limitations** is the field a reader of the artifact would most like to have and is the one an analyst under deadline is most tempted to write as “none.” An artifact with no limitations has not been examined.

E.8 A Worked Sign-Off

The example below runs the checklist on an artifact Chapter 12 has already diagnosed in full, so that the worked example discloses nothing that is not already in the chapter. The artifact is **the four views reproduced and repaired in Lab 12.2**, not the six-chart deck they were taken from. That distinction is the example's first lesson: a sign-off covers what was audited, and signing for six views on the strength of four is precisely the overreach check 2.15 and check 6.5 exist to

catch. The two views Lab 12.2 does not reproduce are a separate artifact and need their own record before anything is claimed about them.

It is written as the record would look after Lab 12.2 Part B, with the repair notes carried across from Section E.6. Its counts are worth reading closely, because the short form for this artifact class is E.5.1 run four times, and the arithmetic below is what makes the sign-off checkable.

Table E.25

A completed sign-off record for the four views of Lab 12.2

Field	Entry
Artifact title	Suburban store performance — the four views reproduced and repaired in Lab 12.2
File and version	suburban_deck_v1, four Tableau worksheets, delivered as PNG exports. Two further views travelled in the same email and are not covered by this record
Data refresh date	Certified extract, line grain and order grain. No refresh date appeared on any of the four as delivered
Population and window	Suburban and resort store types. Two most recent quarters on view one; full certified window on views two through four. Neither the population nor the window was stated on any of the four
Instrument run	Baseline audit: the applicable master checkpoints run in full on each of the four views. Front-end form: E.5.1, run once per view
Completed checklist	appendix_e_signoff_suburban_deck_v1.xlsx, tabs V1 through V4
Checkpoints run	E.1, E.2, E.3, E.4, E.5, E.6, E.7, E.11, E.13, E.14. E.8 not applicable: separate views, no canvas. E.9 not applicable: static exports, no controls. E.10 not applicable: no forecast, estimate, or threshold. E.12 not applicable: working exchange between colleagues, no reader test required
Unique applicable check IDs	89, after the four not-applicable checkpoints and the per-view exclusions recorded on each tab. Counted once for the artifact; the 4 views produce 274 individual check instances, which live on tabs V1 through V4
Unique check IDs passed	71
Unique Level 2 check IDs accepted	1
Unique Level 1 check IDs repaired	17
Unique Level 1 check IDs unresolved	0

Field	Entry
Stop-ship failures found	1.1 and 1.3 — no decision sentence, and the views’ job never declared before they were forwarded as decision support. 2.1, 2.2, and 2.3 — source, window, and filters stated on none of the four. 2.4 and 2.6 — a per-order claim carried on line-grain marks, unverified by hand. 2.12 — no refresh date. 4.3 — the pie puts a magnitude the reader must rank onto angle. 5.1 — view one’s axis runs from \$240,000, drawing the second bar at roughly twelve times the first against a value ratio of 1.046 to one. 5.8 — view two’s dual axis carries two independently scaled ranges. 6.10 — no footer on any of the four. 6.11 — the covering email calls view four “average order value” while the worksheet computes an average of line revenue. 7.1 and 7.2 — a red-to-green palette on an ordered field, with hue carrying the distinction alone. 14.2 and 14.6 — the delivered exports carry neither footers nor a refresh date. Seventeen distinct checks, failing across the four views on 31 instances in all; every instance repaired before recirculation, with its repair note on tabs V1 through V4
Communication revisions	3.9 — the pie replaced by a sorted bar and then by a 100 percent stacked bar split by metro. 3.6 — sort applied to the rebuilt bar. 7.8 — accent hue reserved for the series under discussion on the rebuilt trend view
Accepted defects	11.11 accepted as Caution: no list of charts built and rejected was kept, because the views were drafted in one pass and no alternatives were built. Recorded rather than repaired, and it is the reason the next exchange starts with a specification
Sentence count	Five claim-making sentences checked against 6.2 and 6.4 — four titles and the covering email. Two changed; the email’s “average order value” sentence withdrawn
Removal list	Nothing removed. The pie was replaced rather than cut, and every repair changed an encoding rather than the evidence on display
Usability evidence	Not applicable; a working exchange between colleagues, and the reader test of Checkpoint E.12 is reserved for decision artifacts
AI-use record	Appendix D record for the original request, the response, and the four repairs. Checkpoint E.13: 13.1 failed on the original exchange, since no faults were predicted before the specification was built, and passed on the graded re-run; 13.2 passed, the specification having been built as written with its faults intact; 13.5 failed on four inherited defaults; 13.12 and 13.13 completed
Remaining limitations	The repaired views establish that suburban and resort revenue rose 4.6 percent quarter over quarter. They do not establish that the movement is larger than ordinary quarter-to-quarter variation in this series, and the eight-quarter trend view built in the repair is the artifact that would answer it. Nothing here covers the two views that travelled in the same email
Analyst of record	[Name], [date]
Reviewer	[Name], [date]

Two things about that record are worth naming. The first is arithmetic: 71 passed plus 1 Caution plus 17 Stops repaired plus 0 unresolved is 89, which is the unique-applicable figure, and a reader who cannot make those five numbers reconcile has found a defect in the record rather than in the artifact. That is the whole purpose of the count line, and it is a check anyone can run in ten seconds without seeing the views. Note what the line counts and what it does not: 6.10 failed on all four views and appears once here as one identifier, with four instances and four repairs recorded on the tabs.

The second is why the record is kept at all. The review cost roughly thirty minutes against material that had already reached the director of real estate, and what it produced of lasting value is not the four repaired views — those will be superseded within a week — but the line in the limitations field, which says what the corrected evidence does and does not establish, and the Caution on 11.11, which says what the next exchange should start with. Neither sentence existed before the checklist was run.

E.9 Checkpoint Index by Chapter Section

Table E.26 is the reverse lookup. A student working through a chapter section can find the checkpoint that enforces it without reading this appendix from the front, and an instructor can confirm that a chapter's obligations are represented.

Table E.26

Chapter sections and the checkpoints that enforce them

Chapter section	What it establishes	Checkpoint
2.7; 12.9	The baseline discipline, and its visual form as the reference line	E.5
3.3; 3.4; 12.10	Grain, levels of measurement, and aggregation in the view	E.2, E.7
12.1; 12.2	Visualization's two jobs, and the drift from the first to the second	E.1
12.3; 12.4	Marks, encodings, and the perceptual hierarchy	E.4
12.5	Chart selection by analytic purpose	E.3
12.6; 12.7	Scales, the zero-baseline obligation by encoding, and the five mechanics	E.5
12.8	Color as an encoding system, and accessibility as correctness	E.7
12.9	Reference lines, annotation, and the footer discipline	E.5, E.6, E.14
12.11	Faceting and small multiples	E.3, E.4
12.12	The AI chart-drafting routine and its nine failure modes	E.13

Chapter section	What it establishes	Checkpoint
12.16	The disclosure standard and the duty to inspect	E.11
13.2	Audience analysis and the three-second test	E.1, E.12
13.3	Exploratory and explanatory artifacts, and choosing the form	E.1
13.4; 13.5	The KPI hierarchy and its expression in layout	E.8
13.6	The assertion title	E.6
13.6	Narrative order across a set of views	E.8
13.7	Preattentive emphasis and data-ink decluttering	E.7, E.8
13.8	Interactivity that communicates and interactivity that conceals	E.9
13.5; 13.12.4	Dashboard sizing, containers, the title band, and alternative text	E.8, E.14
13.9	Uncertainty in decision units, and the four-part recommendation	E.10
13.10	The memo, the usability test, and the critique protocol	E.12
13.11	The AI communication routine and its eight failure modes	E.13
13.15	Emphasis, distortion, and the counterfactual test	E.11
Table 8.5 (three transferable points)	Frame, baseline, and error in decision units	E.1, E.5, E.10
Table 12.5 (V1–V3)	Grain and aggregation; scale and baseline honesty; encoding fits the purpose	E.2, E.5, E.3 and E.4
Table 13.6 (C1–C3)	Every sentence earned; the reader’s takeaway; what was removed is disclosed	E.6, E.12, E.11

E.10 Exclusion Register

A checklist that silently omits an obligation looks identical to one that never considered it. Table E.27 records what this appendix does not contain and why, because several of the entries are decisions rather than oversights.

Table E.27

What this appendix does not contain, and why

Not included	Reason
Chart-selection theory, the perceptual hierarchy's evidence, and the grammar of a chart	Chapter 12 owns all three. Checkpoints E.3 and E.4 enforce the conclusions and cite the sections; an appendix that re-argued them would be a second chapter and would drift from the first.
Upstream data-quality checks	Chapter 4 owns the verification log and the defect catalog. Checkpoint E.2 checks that the artifact is consistent with a certified source; it does not certify the source.
Metric definitions	Appendix B is the canonical dictionary. Checkpoint E.6 checks that a definition accompanies a number a reader could misread; it does not supply the definition.
Statistical guidance on choosing an interval	Chapters 10 and 11 own the intervals. Checkpoint E.10 checks that the interval that exists is displayed honestly and converted correctly.
Prompts for AI assistance	Appendix C owns the templates. Checkpoint E.13 is the audit that runs on their output, and Templates C.16 through C.18 route back here for the repair.
A pass/fail score, a weighting, or a grade	The checklist distinguishes hard errors from judgment calls and records both. Converting that into a single number would suppress the distinction, which is the appendix's most useful property.
Project rubrics	The rubrics are a separate instrument and are released with the project briefs. This appendix is incorporated into Project #2's submission workflow as the sign-off record of Section E.7, whatever the rubrics' final home.
Screenshots of software interfaces	Numbered checks and build paths in the chapters are more durable and more accessible than an image of a menu that will be redrawn.
Vendor-specific rules	Every check is written to be tool-neutral. Where a Tableau behavior is named — reference bands, filter scope, alt text, action matching — it is because the behavior is what makes the check necessary, and the vendor's documentation governs the current interface.

One dependency is recorded rather than excluded. Checkpoint E.13 requires an Appendix D record for every AI-assisted exchange, and Chapters 12 and 13 both make the same requirement at the point where the exchange occurs. Appendix D holds the canonical field set — the exchange identifier, the stage, the tool and version, the prompt as sent, the output received, the disposition, the error or limitation found, the verification performed, the change to the final work, and the analyst's own judgment — and Checkpoint E.13 uses those fields rather than defining its own. Where an artifact's AI record and its visual sign-off are submitted together, the

exchange identifier is what links a repair note in Section E.6 to the exchange that made it necessary.

References

- Cleveland, W. S., & McGill, R. (1984). Graphical perception: Theory, experimentation, and application to the development of graphical methods. *Journal of the American Statistical Association*, 79(387), 531–554. <https://doi.org/10.1080/01621459.1984.10478080>
- Correll, M., Bertini, E., & Franconeri, S. (2020). Truncating the y-axis: Threat or menace? In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems* (pp. 1–12). Association for Computing Machinery. <https://doi.org/10.1145/3313831.3376222>
- 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., & Bostock, M. (2010). Crowdsourcing graphical perception: Using Mechanical Turk to assess visualization design. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems* (pp. 203–212). Association for Computing Machinery. <https://doi.org/10.1145/1753326.1753357>
- 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., & 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.
- Krug, S. (2014). *Don't make me think, revisited: A common sense approach to web usability* (3rd ed.). New Riders.

- Munzner, T. (2014). *Visualization analysis and design*. CRC Press.
- 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>
- Pandey, A. V., Rall, K., Satterthwaite, M. L., Nov, O., & Bertini, E. (2015). How deceptive are deceptive visualizations? An empirical analysis of common distortion techniques. In *Proceedings of the 33rd Annual ACM Conference on Human Factors in Computing Systems* (pp. 1469–1478). Association for Computing Machinery. <https://doi.org/10.1145/2702123.2702608>
- 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). *Actions*. Tableau Help. Retrieved August 2, 2026, from <https://help.tableau.com/current/pro/desktop/en-us/actions.htm>
- Tableau. (2026b). *Author views for accessibility*. Tableau Help. Retrieved August 2, 2026, from https://help.tableau.com/current/pro/desktop/en-us/accessibility_create_view.htm
- Tableau. (2026c). *Filter data from your views*. Tableau Help. Retrieved August 2, 2026, from <https://help.tableau.com/current/pro/desktop/en-us/filtering.htm>
- Tableau. (2026d). *Reference lines, bands, distributions, and boxes*. Tableau Help. Retrieved August 2, 2026, from https://help.tableau.com/current/pro/desktop/en-us/reference_lines.htm
- Tableau. (2026e). *Size and lay out your dashboard*. Tableau Help. Retrieved August 2, 2026, from https://help.tableau.com/current/pro/desktop/en-us/dashboards_organize_floatingandtiled.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>
- W3C. (2024). *Web content accessibility guidelines (WCAG) 2.2* (W3C Recommendation, December 12, 2024). World Wide Web Consortium. <https://www.w3.org/TR/WCAG22/>

Wong, B. (2011). Points of view: Color blindness. *Nature Methods*, 8(6), 441.
<https://doi.org/10.1038/nmeth.1618>