

# The Life Sciences Presentation Report

The Perception Science Behind Clinical Data  
Communication

# Index

|                                                                                   |    |
|-----------------------------------------------------------------------------------|----|
| The Perception Hierarchy: Which Chart Encodings Are Actually Read Accurately..... | 3  |
| The Kaplan–Meier Comprehension Problem .....                                      | 6  |
| Axis Truncation: What the Research Actually Found .....                           | 9  |
| Colour: The Accessibility Failure Hiding in Clinical Figures .....                | 11 |
| The Operational Layer: Review Cycles and Content Utilisation .....                | 14 |
| Editor's View — Vaidehi Shukla .....                                              | 18 |
| A Practical Standard for Clinical Data Slides .....                               | 20 |
| Methodology & Sources .....                                                       | 21 |

# 1. The Perception Hierarchy: Which Chart Encodings Are Actually Read Accurately

## 1.1 The foundational finding

In 1984, William Cleveland and Robert McGill published a set of controlled experiments in the Journal of the American Statistical Association establishing how accurately people extract quantitative values from different visual encodings. It remains the empirical foundation of the field.

Their ranking, most accurate to least:

| Rank | Encoding                          | Example in clinical use             |
|------|-----------------------------------|-------------------------------------|
| 1    | Position along a common scale     | Dot plot, aligned bar chart         |
| 2    | Position along non-aligned scales | Small multiples, panelled subgroups |
| 3    | Length, direction, angle          | Unaligned bars, slope charts        |
| 4    | Area                              | Bubble charts, proportional circles |
| 5    | Volume, curvature                 | 3D effects, curved surfaces         |
| 6    | Shading, colour saturation        | Heatmaps, intensity gradients       |

The measured differences are substantial: **position judgments were found to be 1.4 to 2.5 times more accurate than length judgments, and roughly 1.96 times more accurate than angle judgments.**

[Source:](#) 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).

## FIGURE 1

## Perception accuracy ladder

How accurately different visual encodings are perceived.  
Ranked from most accurate (top) to least accurate (bottom).



### The position-vs-angle gap

We perceive position on a common scale almost **twice as accurately** as angle.

Position (Rank 1) **1.96x** Angle (Rank 3)

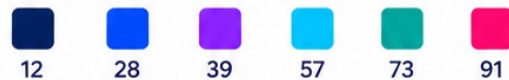


**1.96x more accurate**

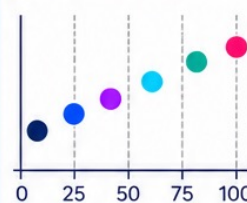
Source: Based on Cleveland & McGill (1984)  
*Graphical Perception: Theory, Experimentation, and Application*

### Worked example

Same dataset (values in %)



Rank 1 — Position  
(most accurate)



Values are encoded as **position** on a common scale.

Rank: 1

Rank 6 — Intensity  
(less accurate)



Values are encoded as **intensity** (lightness).

Rank: 6



### Takeaway

Whenever possible, encode data using **position** on a common scale to maximize perceptual accuracy.

## 1.2 Why this matters for clinical data specifically

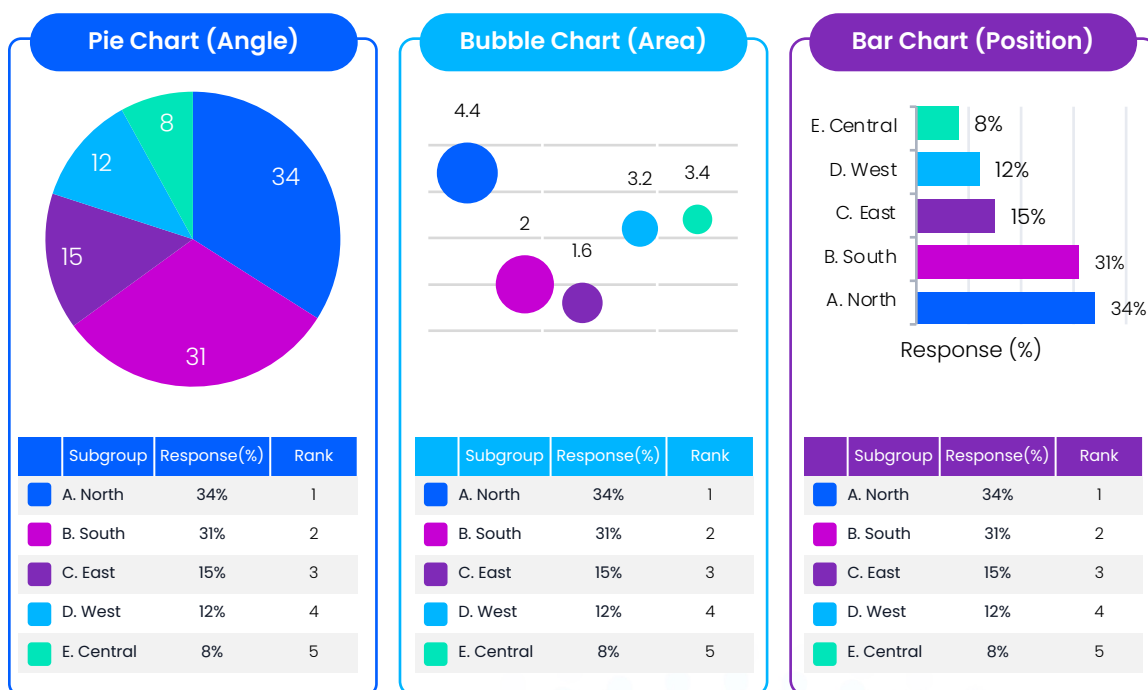
The encodings that rank lowest are exactly the ones clinical and commercial decks reach for most readily: pie charts of patient segments (angle), bubble charts of market opportunity (area), and heatmaps of subgroup response (colour saturation).

A pie chart showing 34% versus 31% response is asking an audience to perform an angle judgment — one of the least accurate tasks available — on a difference that may be the entire point of the slide. The same data as an aligned bar chart is a position judgment, the most accurate task in the hierarchy.

This is not an aesthetic preference. It is a measured difference in how reliably the audience extracts the number.

**FIGURE 2** Three encodings of one dataset

Same clinical response data rendered three ways: angle, area, and position



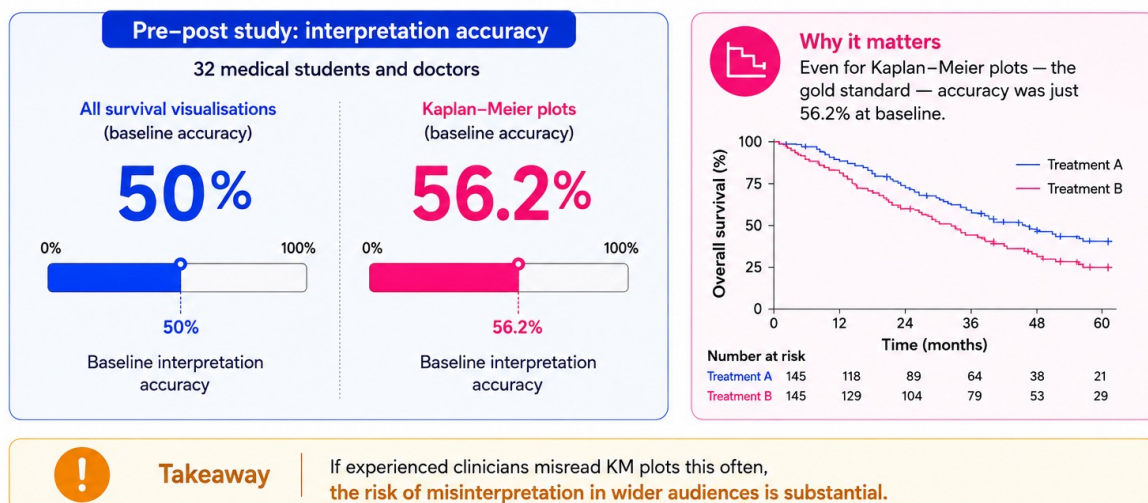
**i** Data (synthetic): Percent of respondents selecting a given option within each subgroup. Total =100%

# 2. The Kaplan–Meier Comprehension Problem

## 2.1 What the measurement shows

The Kaplan–Meier curve is arguably the single most consequential visual in oncology and time-to-event communication. It is also widely misread — including by clinically trained audiences.

A pre-post study of 32 medical students and doctors measured baseline interpretation accuracy for survival visualisations at 50%, with accuracy on Kaplan–Meier plots specifically at 56.2%.



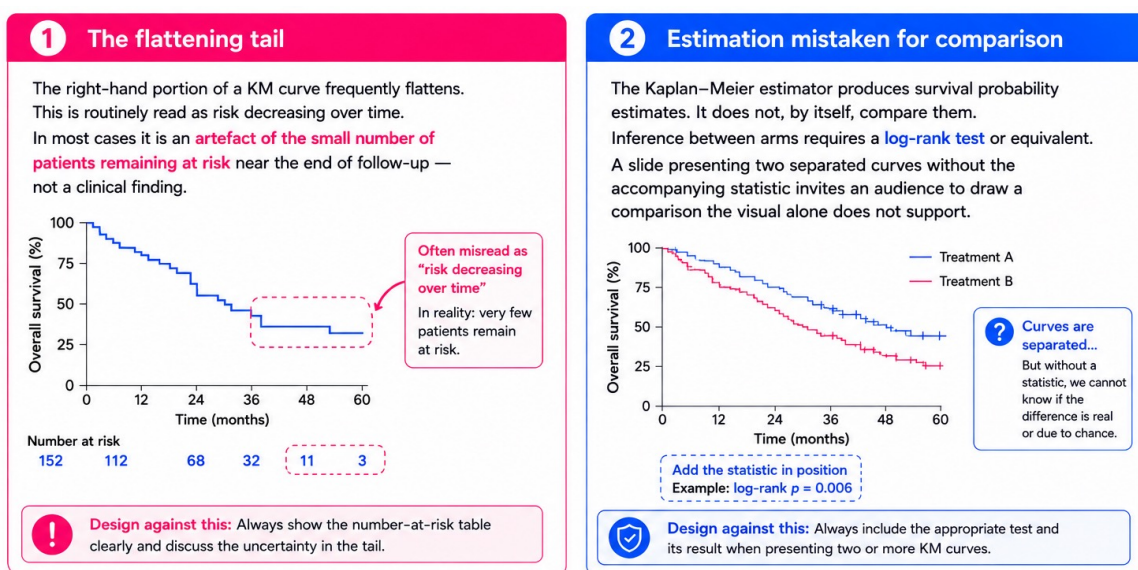
*Source: "Clinicians' Interpretation and Preferences for Survival Data Visualisation: A Pre-Post Study Comparing Kaplan–Meier and Mean Residual Life Plots," arXiv:2511.08332. Note the modest sample (n=32) — this is indicative, not definitive, and should be presented as such.*

## 2.2 The two recurring misreadings

The survival analysis methodology literature identifies two errors consistently enough that they are worth designing against directly.

The flattening tail. The right-hand portion of a KM curve frequently flattens. This is routinely read as risk decreasing over time. In most cases it is an artefact of the small number of patients remaining at risk near the end of follow-up — not a clinical finding.

Estimation mistaken for comparison. The Kaplan–Meier estimator produces survival probability estimates. It does not, by itself, compare them. Inference between arms requires a log-rank test or equivalent. A slide presenting two separated curves without the accompanying statistic invites an audience to draw a comparison the visual alone does not support.



*Sources:* "Kaplan–Meier Survival Analysis: Practical Insights for Clinicians" (PubMed 38631048); "Methods to Analyse Time-to-Event Data: The Kaplan–Meier Survival Curve" (PMC8478547); "The use and abuse of survival analysis and Kaplan–Meier curves in surgical trials," *Journal of Neurological Sciences*.

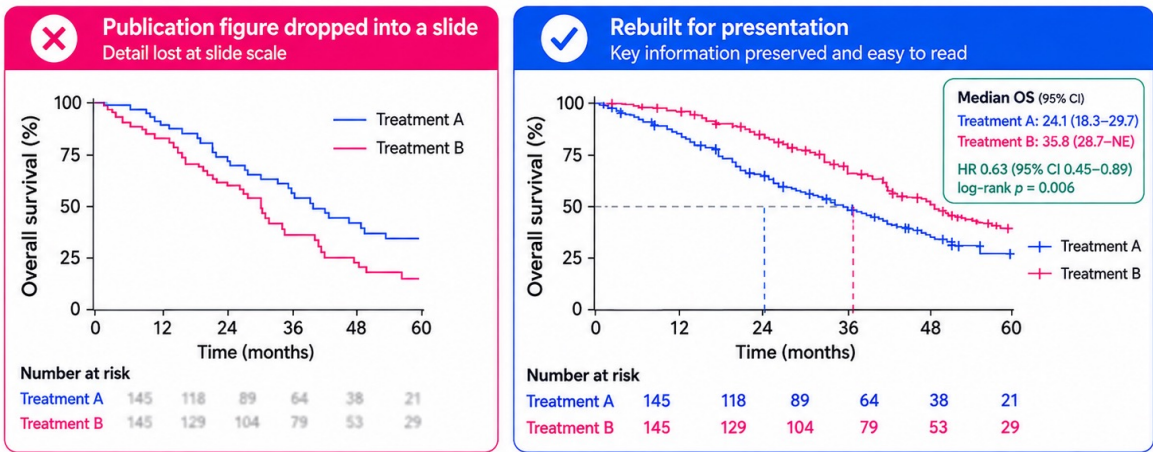
## 2.3 The design implication

Both errors are addressable in the visual itself: number-at-risk tables rendered legibly rather than as a footnote afterthought, censoring marks that remain visible at presentation scale, and the comparative statistic placed with the curves rather than on a following slide.

A KM curve reproduced from a publication at slide scale routinely loses all three — the number-at-risk row becomes unreadable, censoring ticks disappear, and the statistic migrates elsewhere. The chart looks correct and has quietly lost the elements that prevent misreading.

**FIGURE 3** Annotated KM curve: what gets lost at slide scale

Two versions of the same synthetic KM curve. Left: publication figure dropped into a slide — **number-at-risk illegible**, **censoring marks vanished**. Right: rebuilt for presentation — **number-at-risk legible**, **censoring visible**, **statistic in position**.

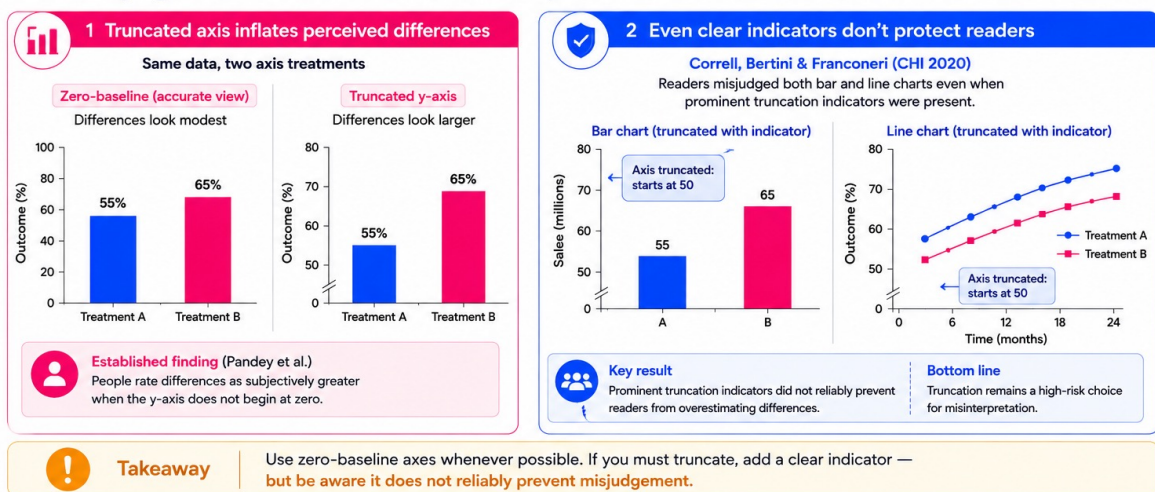


# 3. Axis Truncation: What the Research Actually Found

## 3.1 The measured effect

Truncating a y-axis so it does not begin at zero makes differences appear larger. This is well established: Pandey and colleagues demonstrated that people rate differences in data as subjectively greater in bar charts when the y-axis does not begin at zero.

The more useful finding is more recent and more uncomfortable. Correll, Bertini, and Franconeri (Tableau Research / Northwestern) tested whether clearly signalling a truncated axis protects against the effect. It largely does not – readers misjudged both bar and line charts even when prominent indicators showed the axis had been truncated.



**Sources:** Correll, M., Bertini, E., Franconeri, S., "Truncating the Y-Axis: Threat or Menace?" CHI 2020 (arXiv:1907.02035); Pandey et al., on deceptive visualisation.

## 3.2 The nuance that matters

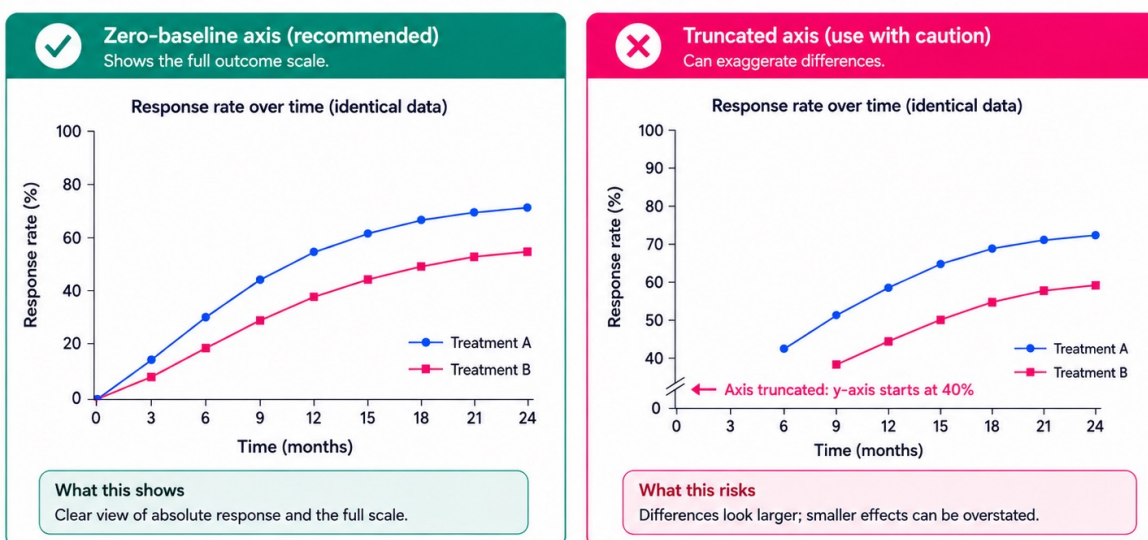
The same research explicitly rejects a blanket prohibition. Correll et al. present cases where truncation is appropriate and cases where it is harmful – the determining factor is communicative and analytic intent, not a universal rule.

For regulated communication this distinction is critical. A truncated axis on an efficacy chart is not automatically non-compliant. But because the research shows disclosure does not neutralise the perceptual effect, a truncated axis on a benefit claim carries a real risk of creating an impression the underlying data does not support – which is a fair balance question, not a design question.

### FIGURE 4 Same data, two axis treatments

Identical synthetic efficacy data. Left: zero-baseline. Right: truncated, with a prominent truncation indicator.

Caption notes that research finds the indicator does not reliably prevent misjudgement.



**Research note** Studies show that even when a truncation indicator is present and prominent, viewers still misjudge the size of differences.



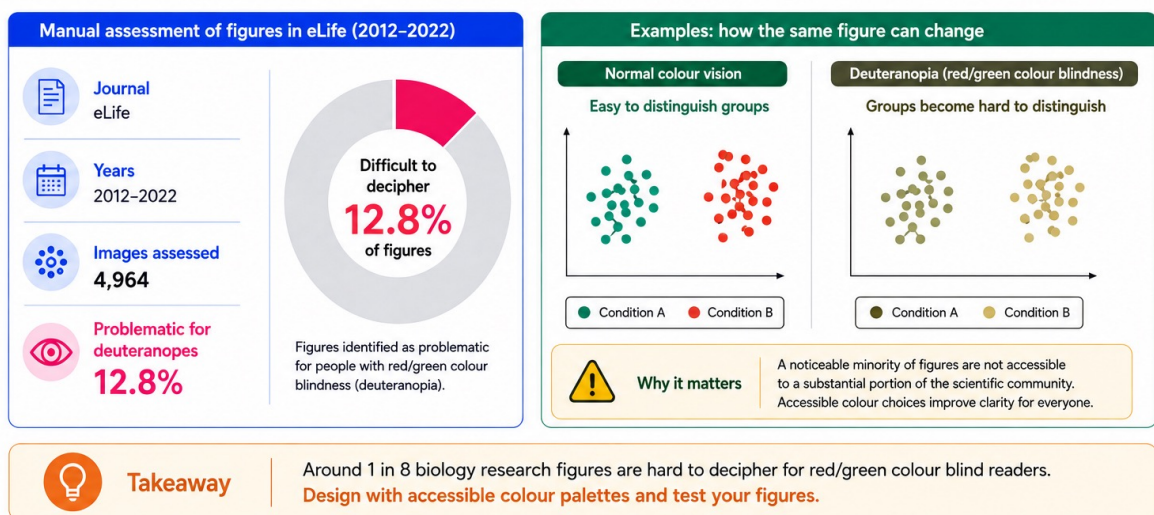
#### Takeaway

Use zero-baseline axes whenever possible. If you must truncate, use a clear indicator — but be aware it **does not reliably prevent misjudgement**.

# 4. Colour: The Accessibility Failure Hiding in Clinical Figures

## 4.1 The scale of the problem

Approximately **13% of figures in biology-related research articles are difficult for people with red/green colour blindness to decipher**. That figure comes from manual assessment of 4,964 randomly selected images published in eLife between 2012 and 2022, where 12.8% were identified as problematic for deuteranopes.



*Source: "Identifying images in the biology literature that are problematic for people with a color-vision deficiency," eLife (2024), article 95524.*

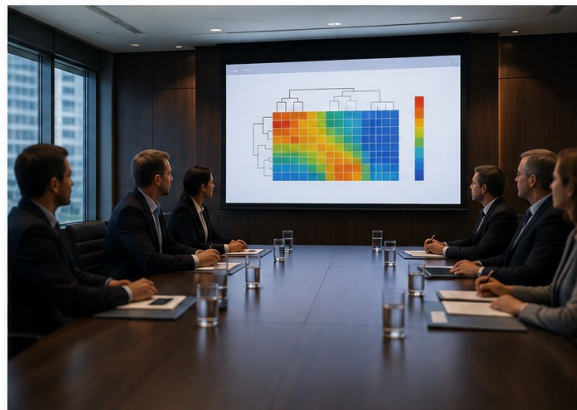
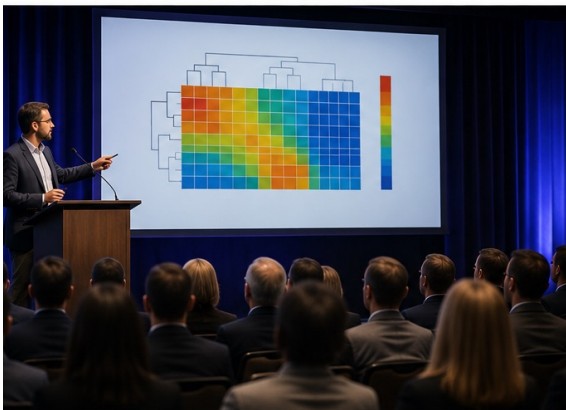
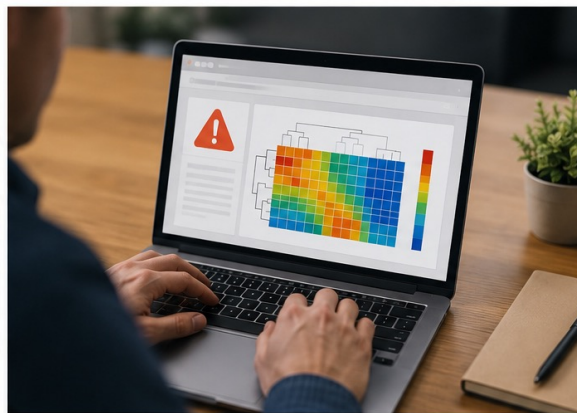
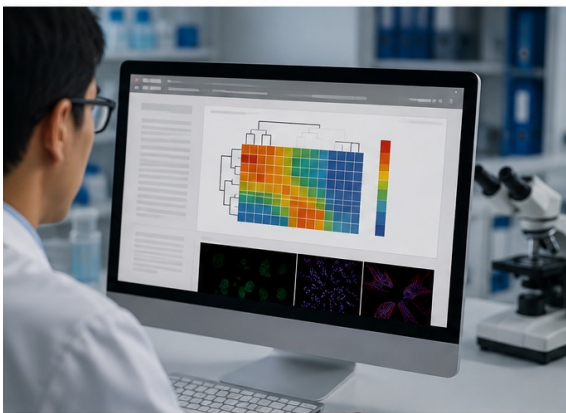
Against an audience prevalence of **red-green colour vision deficiency affecting roughly 7–10% of men**, this is a material communication failure occurring inside peer-reviewed literature – and those same figures are routinely lifted directly into slide decks.



## 4.3 Why this reaches decks intact

Journals have begun deploying automated screening — JetFighter scans bioRxiv preprints for rainbow-based schemes and notifies authors. But figures that predate such screening, or that clear it, are reproduced into congress slides, advisory board decks, and payer submissions without further check.

A figure that was marginal at journal scale, on a printed page, viewed at reading distance, is materially worse projected in a conference room at the back of the audience.



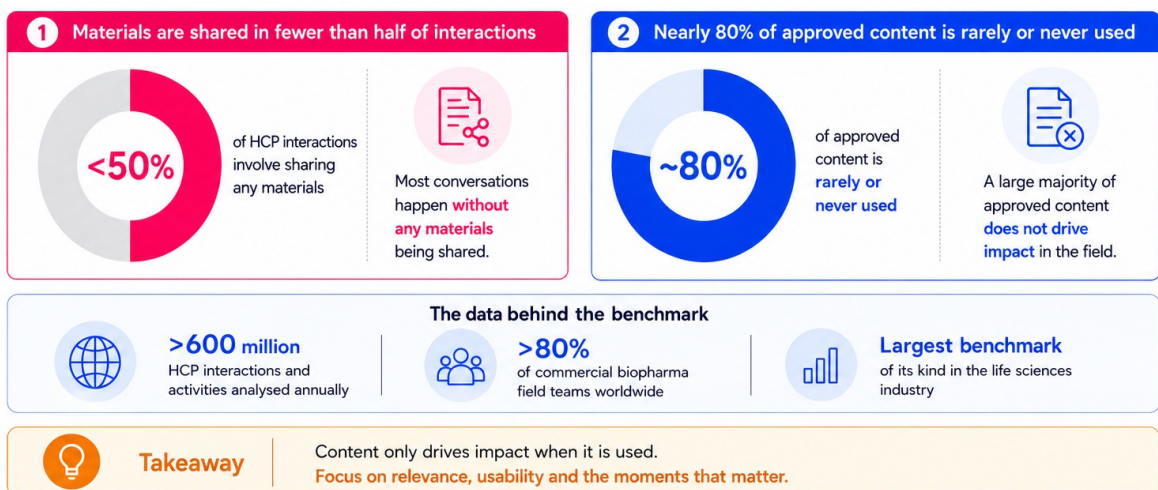
# 5. The Operational Layer: Review Cycles and Content Utilisation

Perception governs whether evidence is understood. Two operational findings govern whether it arrives at all.

## 5.1 Content utilisation

**Field teams share materials in fewer than half of HCP interactions, and nearly 80% of approved content is rarely or never used.**

This comes from the largest benchmark of its kind: Veeva Pulse analyses over **600 million HCP interactions and activities annually, drawn from more than 80% of commercial biopharma field teams worldwide.**



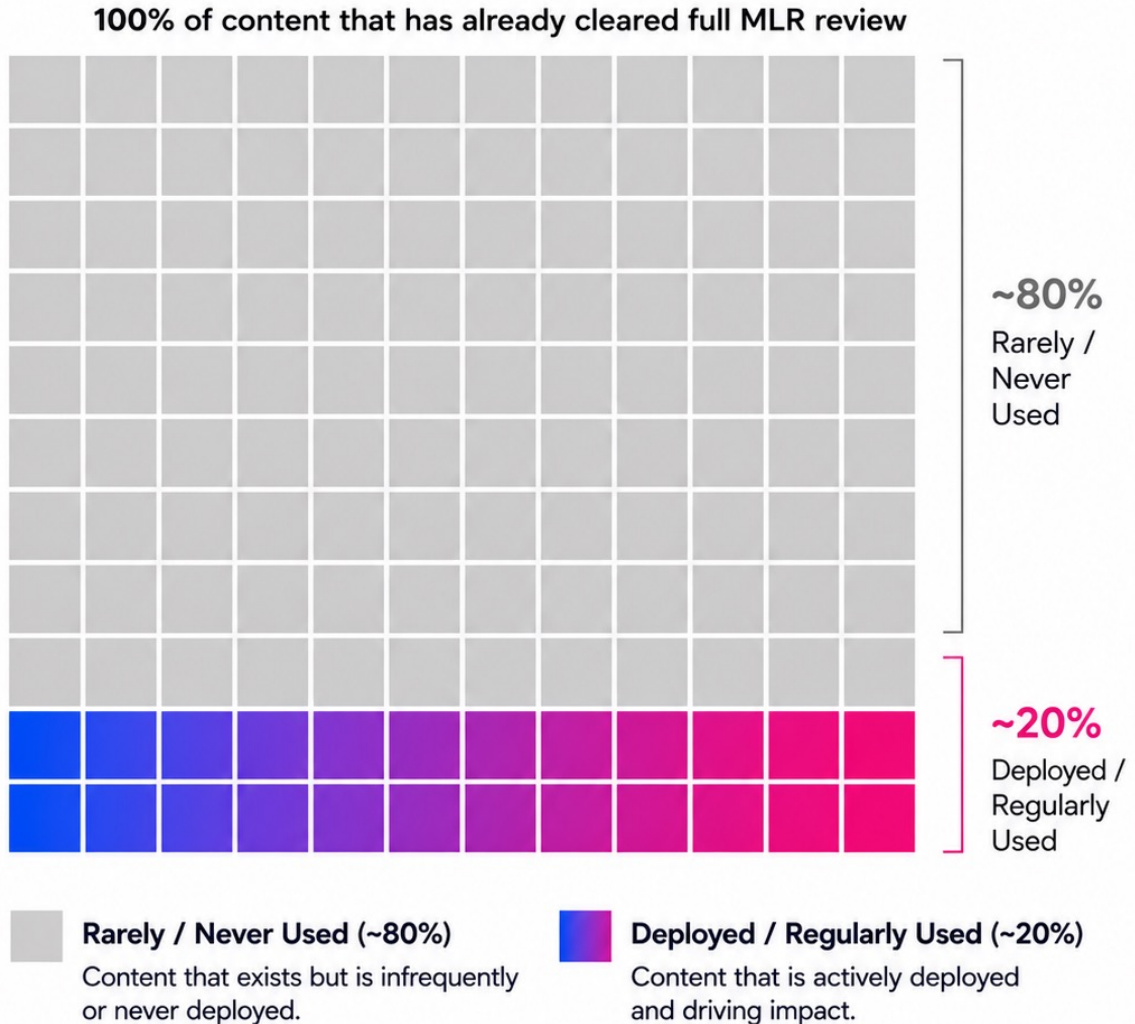
*Source: Veeva Pulse Field Trends Report, published 29 May 2025.*

The cost is not principally production hours. It is the **medical, legal, and regulatory reviewer capacity** consumed approving material that will not be deployed – the scarcest resource in the chain. Every unused asset that cleared review displaced one a field team was waiting for.

## FIGURE 6

# Content utilisation waffle grid

Across typical companies, the vast majority of reviewed content is rarely or never used, while a small proportion drives most impact.



Each square represents **1%** of content  
that has already cleared full MLR review.

## 5.2 Review cycle duration

| Metric                                                                   | Figure                                 | Source                                     |
|--------------------------------------------------------------------------|----------------------------------------|--------------------------------------------|
| MLR review cycle, mid-to-large pharma                                    | 50–60 days per content piece           | Indegene, Pharma Content Compliance Report |
| European market average                                                  | ~20 days, ~1.3 review cycles per asset | Veeva Content Benchmark Report (2023)      |
| Reduction in review cycle time (customer-reported)                       | 57%                                    | Veeva customer data                        |
| Reduction in time spent in MLR/PRC meetings                              | 55%                                    | Veeva customer data                        |
| Compression of regulatory submission timelines via AI-enabled automation | 50–65%                                 | McKinsey (2025)                            |

*Note on the McKinsey figure: it concerns regulatory submission dossiers, not promotional MLR review. The underlying dynamic is comparable but the scope differs, and it should be cited with that distinction intact.*

### FIGURE 7 Review cycle comparison

Documented end-to-end cycle times for comparable review processes.

| PROCESS                                                                                                                                   | TYPICAL CYCLE TIME (CALENDAR DAYS)                                                               | SOURCE                                            |
|-------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------|---------------------------------------------------|
|  <b>Internal MLR review</b><br>(biopharma companies)   |  35–60 days   | Pharma Benchmarks 2023                            |
|  <b>FDA 510(k) review</b><br>(medical devices)         |  90–150 days  | FDA Performance Report 2023                       |
|  <b>EMA centralized procedure</b><br>(human medicines) |  210–277 days | EMA Performance Report 2023                       |
|  <b>UK MHRA review</b><br>(innovative medicines)       |  210–300 days | MHRA Corporate Report 2023/24                     |
|  <b>Health Canada review</b><br>(drugs)                |  240–360 days | Health Canada Report on Plans and Priorities 2023 |

 **MLR cycles are typically measured in weeks, while regulatory review cycles are measured in months.**

## 5.3 Where the time is actually lost

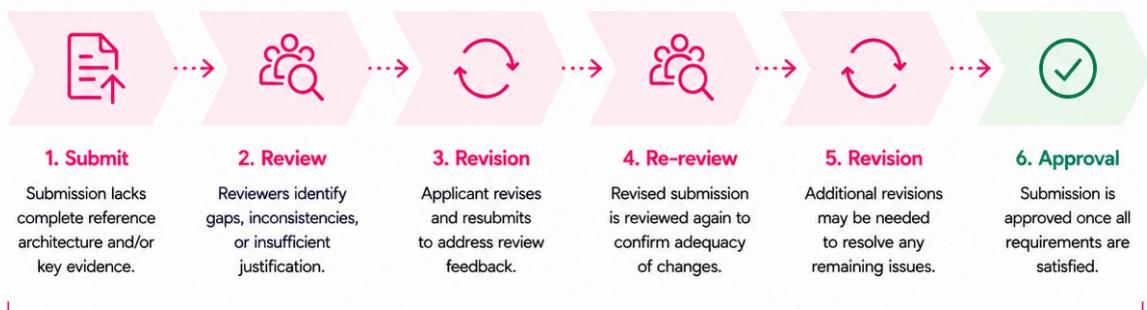
The consistent finding across the operational literature is that delay concentrates before review rather than within it. Incomplete reference packs, claims submitted without linked substantiation, sequential rather than parallel routing, and version fragmentation across email threads all reset cycles.

The corollary: **most of the controllable variance sits with the originating team, not the reviewer.** Content assembled with references linked, claims mapped to approved language, and fair balance addressed at first draft enters review to be confirmed. Content that is not enters review to be corrected — and each correction is a cycle.

### FIGURE 8 Two routes through review

The speed of approval depends heavily on the completeness of the submission.

#### A. Submission without complete reference architecture (longer, iterative path)

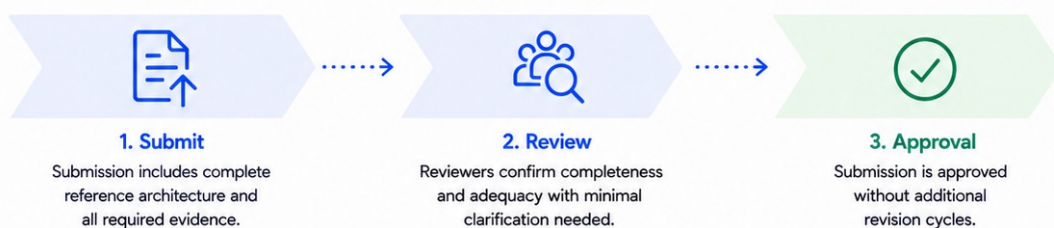


Longer cycle time with multiple iterations



**More iterations = more time.** Incomplete submissions lead to repeated cycles of review and revision, extending time to approval.

#### B. Submission with complete reference architecture (streamlined path)



Shorter cycle time with fewer handoffs



**Both routes lead to the same destination.**  
Complete, well-structured submissions get there faster.

## 6. Editor's View — Vaidehi Shukla

In my experience spanning almost two decades, I see a clear market gap, companies/organisations invest heavily in analysis and generating robust evidence, but comparatively less attention is paid to whether that evidence can be interpreted accurately and consistently by the intended audience. There is a difference in the depth of understanding between a clinician, an MSL and a field representative.

**1. Perception hierarchy:** scientific accuracy does not guarantee understanding: The finding that clearly demarcates evidence interpretation is the visual representation. Both teams are trained to scrutinise endpoints, patient populations, statistical significance, comparators and data sources. The perception research challenges this assumption: values presented through position on a common scale can be interpreted substantially more accurately than those communicated through angles, areas or colour intensity. A pie chart or bubble chart may therefore be numerically correct while still making the intended comparison unnecessarily difficult to decode.

The implication is that visual encoding should be treated as part of evidence quality, not as a downstream design preference. When a decision depends on the audience recognising a precise difference, aligned bars or dot plots should generally take precedence over pie charts, bubbles or heatmaps. This matters particularly in advisory boards, payer discussions and senior leadership meetings, where complex evidence is reviewed under time pressure. Clearer presentation does not dilute the science; it protects its decision-making value by reducing the risk that a valid insight is misunderstood, overlooked or interpreted differently across stakeholders

**2. Kaplan–Meier curves:** rebuild for the decision context, rather than reproduce from the publication: When Kaplan–Meier curves are transferred from a publication into a presentation, they are most commonly pasted, cropped and reduced to fit the available slide space. The curves themselves remain visible, but the elements that protect against misinterpretation are often the first to disappear: the number-at-risk table becomes illegible, censoring marks are lost, endpoint and population definitions move into a small footnote, and the hazard ratio or comparative statistic is separated from the visual. The chart retains the appearance of scientific authority while losing much of the context required to interpret it responsibly. My practical recommendation is to rebuild Kaplan–Meier curves for presentation rather than simply resize them.

The slide should clearly retain aspects like the analysis population, sample size, data cut-off, follow-up period, number-at-risk row, censoring marks and the relevant comparative statistics, including confidence intervals. This is especially important when the tail of a curve appears to plateau: without visibility on how few patients remain under observation, an audience may interpret statistical instability as durable clinical benefit.

Similarly, visibly separated curves can be mistaken for a statistically validated treatment difference when no appropriate comparison has been performed. Rebuilding the chart is therefore not a cosmetic exercise; it preserves the evidence trail and reduces the risk that ambiguity becomes an unsupported clinical claim

**3. Axis truncation:** disclosure alone does not neutralise the claim risk. The most consequential finding on axis truncation is that clearly identifying a truncated axis does not necessarily prevent the audience from overestimating the difference displayed. In a regulated environment, an axis-break symbol may provide formal disclosure, but it does not automatically correct the perceptual impression created by the chart. Teams should therefore treat axis selection as a claims and fair-balance decision, not simply as a formatting choice delegated to the presentation team.

Example, for bar charts intended to communicate comparative magnitude, a zero baseline should generally remain the default. There may be valid analytical reasons to use a narrower range, like when examining movement within a clinically established physiological range or presenting a detailed longitudinal trend. In such cases, the chart should also show absolute values, relevant clinical thresholds and sufficient contextual range to prevent the difference from appearing disproportionate. A useful review question is: “Would a reasonable reader infer a larger clinical difference from this visual than the endpoint results and statistical evidence support?” If the answer is yes, the chart should be redesigned. Transparency is necessary, but in regulated communication it is not sufficient; the visual must also preserve proportionality.

# 7. A Practical Standard for Clinical Data Slides

Drawn directly from the evidence above. Each item traces to a specific finding rather than to design convention.

**01** **Encode for accuracy, not variety.** Where a value must be read precisely, use position on a common scale. Reserve area, angle, and colour saturation for cases where approximate magnitude is genuinely sufficient. (Cleveland & McGill)

**02** **Rebuild Kaplan-Meier curves for presentation;** do not paste them. Number-at-risk legible at projection scale, censoring marks visible, comparative statistic co-located with the curves. (Survival analysis methodology literature)

**03** **Treat axis truncation as a claims decision, not a formatting decision.** Disclosure does not neutralise the perceptual effect, so the question is whether the impression created is one the data supports. (Correll et al.)

**04** **Test every clinical figure under deuteranopia simulation before it ships.** Roughly 13% of published biology figures fail this test; slides inherit them unchanged. (eLife, 2024)

**05** **Retire rainbow and jet colormaps from clinical data.** They distort values by a measurable margin and collapse under common colour vision deficiencies. (HESS, 2021; RSNA)

**06** **Resolve compliance architecture before submission, not during review.** The controllable variance in cycle time sits upstream of the reviewer. (Indegene; Veeva)

**07** **Build modular, not bespoke.** With nearly 80% of approved content going unused, one review cycle producing many compliant configurations is a materially better use of reviewer capacity than many cycles producing many single-use assets. (Veeva Pulse)

# 8. Methodology & Sources

This report synthesises peer-reviewed research in graphical perception, survival analysis methodology, and scientific visualisation accessibility, together with published industry benchmark data on content operations. No primary survey was conducted. Where a finding rests on a small sample, that limitation is stated in the text.

Sources are grouped by evidentiary weight so readers can assess each on its own terms.

## Peer-reviewed and academic

- 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.
- Correll, M., Bertini, E., & Franconeri, S. (2020). "Truncating the Y-Axis: Threat or Menace?" CHI 2020. <https://arxiv.org/pdf/1907.02035>
- "Identifying images in the biology literature that are problematic for people with a color-vision deficiency." *eLife* (2024), 95524. <https://elifesciences.org/articles/95524>
- "Rainbow color map distorts and misleads research in hydrology – guidance for better visualizations and science communication." *Hydrology and Earth System Sciences*, 25:4549 (2021). <https://hess.copernicus.org/articles/25/4549/2021/>
- "The misuse of colour in science communication." *Nature Communications*. <https://pubmed.ncbi.nlm.nih.gov/33116149/>
- "Kaplan-Meier Survival Analysis: Practical Insights for Clinicians." <https://pubmed.ncbi.nlm.nih.gov/38631048/>
- "Methods to Analyse Time-to-Event Data: The Kaplan-Meier Survival Curve." <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC8478547/>
- "Clinicians' Interpretation and Preferences for Survival Data Visualisation." <https://arxiv.org/pdf/2511.08332> (n=32; indicative)
- Radiological Society of North America (2023). "Scientifically Derived Color Maps and How to Use Them." [https://pubs.rsna.org/page/radiology/blog/2023/5/ryblog\\_05192023](https://pubs.rsna.org/page/radiology/blog/2023/5/ryblog_05192023)

## Industry benchmark data

*Large-sample commercial research. Attributed by name throughout; readers should weigh the commercial context.*

- Veeva Pulse Field Trends Report (29 May 2025) – 600M+ HCP interactions annually, 80%+ of commercial biopharma field teams.  
<https://www.veeva.com/resources/veeva-pulse-field-trends-report-1q25/>
- Veeva Content Benchmark Report (2023)
- Indegene, Pharma Content Compliance Report. <https://www.indegene.com/what-we-think/blogs/mlr-bottlenecks-in-pharma>
- McKinsey (2025), "Rewiring pharma's regulatory submissions with AI and zero-based design"



## For enterprise presentation consulting

[emarketing@aislides.com](mailto:emarketing@aislides.com)

**Learn more:** [www.aislides.com](http://www.aislides.com)

### Important Disclaimer and Terms of Use

This report has been prepared by **AI Slides** for informational and illustrative purposes only. While every effort has been made to ensure the accuracy and completeness of the information presented herein, AI Slides makes no representations or warranties of any kind, express or implied, about the completeness, accuracy, reliability, suitability, or availability with respect to the information, products, services, or related graphics contained in this report for any purpose.

**Case Studies and Performance Data:** The case studies and performance data presented in this report, including but not limited to metrics, outcomes, and client descriptions, are based on real projects undertaken by AI Slides. Certain client details and specific figures have been anonymized or generalized for confidentiality and illustrative purposes. These examples are provided to demonstrate our capabilities and methodology. Past performance is not indicative of future results, and actual outcomes may vary depending on numerous factors specific to each client, project, industry, and market conditions. AI Slides does not guarantee that similar results will be achieved in any future engagement.

**Forward-Looking Statements:** This report may contain forward-looking statements or projections regarding future trends, market conditions, or potential outcomes. These statements are based on current expectations, assumptions, and analyses, and are subject to risks and uncertainties that could cause actual results to differ materially. AI Slides undertakes no obligation to update or revise any forward-looking statements.

**Proprietary Information:** All content within this report, including text, graphics, methodologies (e.g., Insight-First Design), and concepts, is the proprietary property of AI Slides unless otherwise noted, and is protected by copyright and intellectual property laws. Unauthorized reproduction or distribution of this material, in whole or in part, is strictly prohibited.

**No Legal, Financial, or Professional Advice:** The information in this report is not intended to constitute, and should not be relied upon as, professional advice of any kind, including legal, financial, or strategic consulting advice. Readers should consult with their own qualified advisors for specific guidance tailored to their individual circumstances.

**Limitation of Liability:** To the fullest extent permitted by law, AI Slides, its affiliates, directors, employees, and agents shall not be liable for any direct, indirect, incidental, special, consequential, or punitive damages, or any loss of profits or revenues, whether incurred directly or indirectly, or any loss of data, use, goodwill, or other intangible losses, resulting from (a) your access to or use of or inability to access or use the report; (b) any conduct or content of any third party in the report; or (c) unauthorized access, use or alteration of your transmissions or content.

By accessing and reading this report, you acknowledge and agree to these terms.

