

Reporting Results Clearly

Companion to the Week 10 episode. Say everything you found, accurately — without the selective "spin" that distorts the literature.

1. The one idea

Reporting = **completeness** + **calibration**. Selectivity (report the winners, bury the nulls, use stronger-than-warranted language) each feels like "a clear story" but collectively biases the literature toward inflated positives.

2. Report ALL preregistered outcomes

Report every preregistered outcome, incl. nulls — the antidote to **selective outcome reporting** / **outcome switching**. If you tested 5 things and report only the 1 that "worked", it's likely chance — but the reader can't tell. Preregistration lets a reader **check your paper against your plan**, and lets you present a planned null as the *real, publishable finding* it is.

3. Use reporting-guideline checklists

Study type	Guideline
Randomised controlled trial	CONSORT (Schulz, Altman & Moher, 2010) — randomisation, flow of participants, outcomes, effect sizes+CI, dropouts
Observational (cohort/case-control/cross-sectional)	STROBE (von Elm et al., 2007) — 22-item checklist

Work the relevant checklist as you write — it catches what memory misses and makes selectivity harder. (The EQUATOR network collects guidelines for every study type; many journals require them.)

4. Effect sizes & nulls

Make **effect sizes** + **confidence intervals** (interpreted vs your SESOI) the centrepiece — not bare p-values. Report **nulls fully and unapologetically**: each is a small correction to publication bias (Week 1). A well-designed "no meaningful effect" is a genuine contribution.

5. Spin — the quiet distortion

Spin move	Honest, calibrated version
Primary null → spotlight a significant secondary	Abstract conclusion matches the primary result
Tiny effect called "substantial / promising"	Adjectives match the effect size ("small")
Observational data → "X improves Y"	"X was associated with Y" (reserve causal verbs for randomised)
Limitations buried at the end	Limitations plain and early; framed as one contribution

Self-test: read your *abstract alone* — "if this were all someone saw, would they believe something *stronger* than my data support?" If yes, remove the spin.

6. Do this — this week

Completeness-and-calibration audit

1. Every preregistered outcome reported, incl. nulls (cross-check vs preregistration).
2. Effect size + CI everywhere (vs SESOI), not just p.
3. Work the CONSORT/STROBE checklist; fill gaps.
4. Spin check: read the abstract + conclusion alone — is the language, emphasis & causal claim calibrated to the evidence?

References

Schulz, K. F., Altman, D. G., & Moher, D. (2010). CONSORT 2010 statement. *BMJ*, *340*, c332.

von Elm, E., Altman, D. G., Egger, M., Pocock, S. J., Gøtzsche, P. C., & Vandenbroucke, J. P. (2007). The STROBE statement. *Annals of Internal Medicine*, *147*(8), 573–577.

Munafò, M. R., et al. (2017). A manifesto for reproducible science. *Nature Human Behaviour*, *1*, 0021.