



Error & confidence cheat sheet – for claims

The maths behind "how sure are we?" — written for the EIS/yeast work, applicable everywhere. **Rule zero: every number in a claim carries three companions — *n*, an uncertainty, and the conditions it was measured under.**

Bluefruit Software — internal reference

1 Vocabulary

- Accuracy** — closeness to the true value (needs a reference). **Precision** — closeness of repeats to each other. You can be precise and wrong.
- Repeatability** — same operator, rig, day. **Reproducibility** — different operator/day/rig. Quote which one.
- Random error** averages away with more runs. **Bias** doesn't — only a reference reveals it. More *n* never fixes bias.
- SD** — spread of runs. **SEM** = $s/\sqrt{n}$ — uncertainty of the mean. Label which on every plot.

2 The basics

CV = $s/\bar{x}$ (report as %) — the workhorse for repeatability rows.

95% CI of the mean = $\bar{x} \pm t \cdot s/\sqrt{n}$, where *t* depends on *n*:

<i>n</i>	3	4	5	6	10	20	30+
<i>t</i>	4.30	3.18	2.78	2.57	2.26	2.09	≈2.0

At ***n*=3** the CI is ±4.3 SEM — triplicates give an *indication*, not a quantified claim.

3 Comparing two conditions

$\Delta = \bar{x}_1 - \bar{x}_2$, **SE(Δ)** = $\sqrt{(s_1^2/n_1 + s_2^2/n_2)}$, $CI \approx \Delta \pm 2 \cdot SE(\Delta)$. If the CI on Δ excludes zero, the separation is real.

Effect size *d* = Δ/s — "separation is *N*× the noise". *d* < 1: heavy overlap; *d* ≈ 2: barely overlapping; *d* ≥ 5: LED-demo territory.

Eyeball rule: non-overlapping 95% CIs → significant. Overlapping → *not necessarily* insignificant — do the sum.

4 How many runs?

Decide *before* the experiment. Per group, 80% power at $p < 0.05$:

$$n \approx 16 / d^2$$

***d* = 2**
n ≈ 4

***d* = 1**
n ≈ 16

***d* = 0.5**
n ≈ 64

Backwards: *n*=5 per condition only reliably detects $d \geq 1.8$. Write the target *d* into the experiment's acceptance criteria.

5 Error budget

Independent sources add **in quadrature**: $\sigma^2 = \sigma_{\text{fixture}}^2 + \sigma_{\text{prep}}^2 + \sigma_{\text{drift}}^2 + \sigma_{\text{electronics}}^2$. Measure each separately; improve the dominant term, ignore the rest.

Multiplying/dividing: **relative** errors add in quadrature.

Ground-truth cap: $\sigma_{\text{claim}} \geq \sigma_{\text{reference}}$. If cell counting agrees to ±10%, no viability claim can be tighter than ±10% — however good the instrument.

6 Repeatability studies

Report **CV per frequency** across the sweep — repeatability is rarely uniform in frequency.

Separate variance components: **total = between-remount + within-mount**. Recipe: *k* remounts × *m* runs each; within = mean of per-mount variances; between = variance of per-mount means – within/*m*.

The bigger component names the enemy — fixture vs electronics/drift.

7 Classification claims

Report **sensitivity** and **specificity** separately — a 90%-accurate classifier on a 90/10 dataset may be useless.

CI on a proportion: $p \pm 1.96 \cdot \sqrt{p(1-p)/n}$ (*n* ≥ ~30; Wilson for smaller).

19/20 correct = 95% accuracy, but the 95% CI is ~75–99%. Claiming "95% ± 5%" needs ***n* ≈ 70+ blind runs**. Small-*n* accuracies are anecdotes wearing percentages.

8 Trend / dose-response

Fit the agreed model: report **slope ± 95% CI**. "Monotonic response" when the slope CI excludes zero.

Stronger claim: means are ordered *and* adjacent CIs on Δ exclude zero.

Plot **every replicate point**, not just means — reviewers and clients' scientists trust scatter.

Key *n* number of runs $\bar{x}$ mean *s* / **SD** standard deviation **SEM** standard error of the mean ($s/\sqrt{n}$) **CV** coefficient of variation ($s/\bar{x}$) **CI** confidence interval *t* Student's *t* multiplier Δ difference of two means **SE(Δ)** standard error of Δ
d effect size (Δ/*s*) σ SD of one error source *p* probability of a false positive **EIS** electrochemical impedance spectroscopy **Rct** charge-transfer resistance **AC** acceptance criteria

9 Honesty rules — the ones that get checked

Significant figures follow uncertainty: **3.2 kΩ ± 0.3 kΩ** — never 3.217 kΩ ± 0.3 kΩ.

State ***n* on every figure**; label SD vs SEM vs CI on every error bar.

"Inconclusive" with small *n* means **underpowered**, not "no effect" — say which.

Declare the analysis **before the data**; reporting the best of many frequencies inflates false positives. Post-hoc finds are new hypotheses, not results.

Blind before you quantify. Unblinded quantified claims cap at Client-demo rung regardless of the stats.