



FORMULAS YOU MUST KNOW

$\hat{\mu} = \bar{x} = \frac{\sum x}{n}$ — the sample mean estimates the population mean μ .

Estimate of the mean:

$E(\bar{x}) = \mu$ and $E(s_{n-1}^2) = \sigma^2$ — right on average, not too big or too small.

Unbiased means:

$SE = \frac{\sigma}{\sqrt{n}}$ (known σ) or $\frac{s}{\sqrt{n}}$ (unknown σ) — the standard deviation of $\bar{x}$.

Standard error:

$$\bar{x} \pm z^* \frac{\sigma}{\sqrt{n}}.$$

σ known (z -interval):

Enter the raw list → choose TInterval (or ZInterval) → set the C-level, e.g. 0.95.

Data mode:

The interval (a, b) , ready to quote, together with $\bar{x}$ and n .

Output:

$\widehat{\sigma^2} = s_{n-1}^2 = \frac{\sum (x - \bar{x})^2}{n - 1}$ — divide by $n - 1$, not n .

Estimate of the variance:

$\bar{X} \sim N\left(\mu, \frac{\sigma^2}{n}\right)$ — exact for a normal

population, approximate for any population when n is large.

Sampling distribution:

$\bar{x} \pm E$ — always centred on the sample mean.

The interval:

$\bar{x} \pm t^* \frac{s}{\sqrt{n}}$, with t^* read at $\nu = n - 1$ degrees of freedom.

σ unknown (t -interval):

Type in $\bar{x}$, s_x (or σ) and n → same routine, same C-level.

Stats mode:

$E = t^* \frac{s}{\sqrt{n}}$ (or $z^* \frac{\sigma}{\sqrt{n}}$) — critical value $\times$ standard error.

Margin of error:

KEY METHODS

- This is the exact opposite of describing a fixed data set (Univariate Statistics), where IB wants σ_x . Rule of thumb: describing data → σ_x ; estimating a population from a sample → s_x .
- Divide by $\sqrt{n}$, never by n . Quadruple the sample and the standard error only halves — precision improves slowly, and that is exactly why big samples are expensive.
- Central Limit Theorem: for $n \geq 30$ the sample mean is approximately normal whatever the population's

EXAM TRAPS TO AVOID

- Dividing by n (the calculator's σ_x) systematically underestimates the population variance. For estimation IB uses the $n - 1$ divisor — the s_x key — so the estimator is unbiased.
- Pick the interval by what you are given: a stated population $\sigma \rightarrow z$; only the sample $s \rightarrow t$. In AI HL exams σ is almost always unknown, so the t -interval is the default choice.
- E multiplies the standard error $s/\sqrt{n}$ — not s on its own. Dropping the $\sqrt{n}$ inflates every limit by a factor of $\sqrt{n}$.

shape — the reason the interval formulas work so widely in real AI applications.

- $t^* > z^*$ for the same confidence (the t -distribution has heavier tails), and $t^* \rightarrow z^*$ as n grows — for large samples the two intervals nearly coincide.
- For a t -interval, enter the sample standard deviation s_x (the $n - 1$ value) in the s.d. field. For a z -interval enter the given population σ . Wrong field, wrong interval.
- If a question wants the margin of error but you built the interval on the GDC, recover it instantly as $E = \frac{b - a}{2}$.
- t^* depends on both the confidence level and $\nu = n - 1$; read it from the GDC's inverse- t or straight from the interval — never guess it.

- Width is the full span, so keep the factor of 2: $w = 2E$. The sample mean sits at the midpoint, never at an end.
- Raising the confidence level widens the interval but does not move its centre $\bar{x}$ — more confidence buys less precision, not a different estimate.
- Always round n up to the next whole number — rounding down leaves the interval wider than the target. And n must be a positive integer.
- It is not "there is a 95% probability that μ lies in this interval". μ is a fixed number; the interval is what is random. Once computed, this particular interval either contains μ or it does not.