

Marketing analyst: study guide

Sales, marketing, customer service, and contact centres · suite apt-294-marketing-analyst · generated 2026-09-15T15:37:51.861Z

Detail	Value
Title	Marketing analyst: study guide
Generated	2026-09-15T15:37:51.861Z
Fixture/version	apt-294-marketing-analyst
Sector	Sales, marketing, customer service, and contact centres
Guide version	v2

Private practice result. Not an official exam certificate, employer decision, hiring signal, admissions decision, or guaranteed outcome. Scores stay on this device unless you export them.

Answer keys and scoring logic stay server-side and are never included in any download or export.

How to use this guide

This file contains the whole study outline for this suite: every skill it draws on, the full lesson for each of those skills, worked examples, practice tips, a glossary, and where each piece of material comes from. Nothing here is a summary of a page you still have to visit.

1. Read the skills section end to end once, without timing yourself.
2. Work the guided practice mode for the suite, using the practice tips as a checklist.
3. Move to the mini-test only when guided practice feels unhurried.
4. Sit the full simulation last, once, in the conditions you expect on the day.

Practice attempts are stored on the device you used, never on an account. Exporting a result is the only way anything leaves that device.

Practice modes and durations

Practice modes for Marketing analyst

Mode	Duration	What it is for
Guided practice	15 minutes	Untimed, with feedback after every item.
Mini-test	18 minutes	A short timed set for checking pace.
Full simulation	45 minutes	Full length and full time, in one sitting.

- Start guided practice:
<https://learn.novusstreamsolutions.com/aptitude/practice/new?bank=apt-294-marketing-analyst&mode=guided>
- Start mini-test:
<https://learn.novusstreamsolutions.com/aptitude/practice/new?bank=apt-294-marketing-analyst&mode=mini>
- Start full simulation:
<https://learn.novusstreamsolutions.com/aptitude/practice/new?bank=apt-294-marketing-analyst&mode=full>

Skills covered, in full

This suite draws on 5 skill constructs. Each one below carries its complete lesson.

Numerical reasoning

Arithmetic, fractions, percentages, ratios, rates, estimation, word problems, and number relationships.

Numerical reasoning is the ability to take a small set of supplied numbers (a price list, a staffing table, a fuel figure), and reach a defensible answer in around ninety seconds without a formula sheet. It is the most widely used quantitative section in graduate, management and public-sector screening, and it appears in banking, retail, armed forces and healthcare entry batteries alike. The arithmetic itself is deliberately ordinary: percentages, ratios, rates and averages. What is actually being measured is whether you pick the right operation quickly, keep the units straight, and answer the question that was asked rather than the one you started computing. Practice here is device-local: no account, and nothing leaves this device unless you export it.

What you should be able to do after this lesson:

1. Move between fractions, decimals, percentages and ratios on sight, and pick whichever form makes a given item fastest (12.5 percent as one eighth, 33.3 percent as one third).
2. Apply percentage change in both directions, including recovering an original figure from a post-increase total, and combine successive changes multiplicatively rather than by adding them.
3. Split a total in a stated ratio by counting parts first, and work backwards from a stated difference between two shares.
4. Solve rate problems (speed, unit price, consumption, combined work rates) by carrying units through every intermediate line so the units of the answer prove the method.
5. Produce a one-significant-figure estimate before computing, and use it to eliminate the distractors built from the common wrong operations.
6. Read a multi-step stem and name, in one sentence, the single quantity the final question asks for before touching the numbers.

Worked examples and pitfalls

Reverse percentages: the item most candidates get backwards: A subscription price rose by 15 percent and now stands at 57.50 pounds. What was it before? The correct move is to divide, not subtract: the new price is 115 percent of the old, so the old price is $57.50 / 1.15 = 50.00$. Check it forwards: 15 percent of 50 is 7.50, and $50 + 7.50 = 57.50$. The tempting wrong answer takes 15 percent of the NEW figure, $0.15 \times 57.50 = 8.625$, and reports 48.88. That distractor is always on the option list because it is what most people do under time pressure, and it is wrong by 2.25 percent, small enough to look plausible. The same asymmetry drives the successive-change item: a price that rises 20 percent and then falls 20 percent does not return to where it started. Multiply the factors: $1.20 \times 0.80 = 0.96$, a net fall of 4 percent. Starting at 500 pounds you get 600 then 480, not 500. Percentages compose by multiplication; they never add.

Ratio splits: count the parts before you divide: A 4,830 pound budget is split between three teams in the ratio 3:5:6. Add the parts first: $3 + 5 + 6 = 14$. One part is $4,830 / 14 = 345$. The shares are $3 \times 345 = 1,035$, $5 \times 345 = 1,725$ and $6 \times 345 = 2,070$, and they sum back to 4,830, which is the check you should always run. Two classic failures. The first is dividing by 3 because there are three teams. That gives 1,610 each and ignores the ratio entirely. The second is treating 5 as a fraction and computing five sixths or five fourteenths of something other than the total. The more interesting version of this item gives you a difference instead of the total: 'the second team receives 690 pounds more than the first, what is the whole budget?' The difference between the shares is $5 - 3 = 2$ parts, so one part is $690 / 2 = 345$, and the budget is $14 \times 345 = 4,830$. Same number, reached from the other end, and the arithmetic is trivial once you have made the parts

explicit.

Rates and units: 2 h 15 min is 2.25, not 2.15: A delivery van covers 174 km in 2 hours 15 minutes. Average speed is distance over time, and the time must be in hours: 15 minutes is $15/60 = 0.25$ h, so $174 / 2.25 = 77.3$ km/h. Type 2.15 into the calculator instead and you get 80.9 km/h: an error of roughly 4.7 percent that sits comfortably inside the plausible range and will match one of the options. Now extend it. The van consumes 8.6 litres per 100 km, so the trip needs $174 \times 8.6 / 100 = 14.964$ litres, and at 1.48 pounds per litre that is $14.964 \times 1.48 = 22.15$ pounds. Notice the units doing the work: (km) $\times$ (L / 100 km) leaves litres; (L) $\times$ (pounds / L) leaves pounds. If your intermediate line has km still attached at the end, you have divided when you should have multiplied. Write the unit next to every number and the method checks itself.

Unit pricing: normalise before you compare: Three pack sizes of the same fluid. Pack A: 750 ml for 3.60 pounds. Pack B: 2 litres for 9.40. Pack C: 1.5 litres for 7.20. Convert all three to price per litre. A: $3.60 / 0.75 = 4.80$ per litre. B: $9.40 / 2 = 4.70$. C: $7.20 / 1.5 = 4.80$. So B is cheapest by 10p a litre, and A and C are identical despite looking like different deals. The 'bigger pack is cheaper' heuristic happens to hold here but is not a rule, and test writers know it. Half of these items are built specifically so the largest pack loses. Now add the multibuy that these questions love: pack A is on three-for-two. Three packs give 2.25 litres for the price of two, 7.20 pounds, which is $7.20 / 2.25 = 3.20$ per litre and beats everything. The trap in the multibuy version is dividing by the number of packs paid for rather than the volume received.

Combined work rates: add rates, never times: Printer A completes a 3,000-page run in 50 minutes; printer B completes the same run in 75 minutes. Running together, how long? Convert to rates: A prints $3,000 / 50 = 60$ pages per minute, B prints $3,000 / 75 = 40$ pages per minute, and together they print 100 pages per minute, so the job takes $3,000 / 100 = 30$ minutes. Verify by counting output: in 30 minutes A produces 1,800 pages and B produces 1,200, which is 3,000 exactly. The two wrong answers you will see on the option list are 62.5 minutes (the average of 50 and 75) and 125 minutes (the sum). Both are impossible on inspection: two machines working together must finish faster than the faster machine alone, so any answer above 50 minutes is wrong before you compute anything. The general form is $1 / (1/50 + 1/75)$, and it is worth being able to write that line directly, but the sanity bound catches the error faster than the algebra does.

Estimate first, then compute: killing distractors in ten seconds: 'A department spent 847,300 pounds in 2024. In 2025 the budget fell by 12.5 percent. What was the 2025 spend?' Options: 741,387.50 / 953,212.50 / 105,912.50 / 762,570. Estimate before anything else: 12.5 percent is exactly one eighth, one eighth of roughly 850,000 is roughly 106,000, so the answer is roughly 744,000. That single line eliminates three options. 953,212.50 applied the change upwards. 105,912.50 is the size of the reduction, not the resulting spend: the 'answered the wrong question' distractor, and the most commonly selected wrong option on items of this shape. 762,570 is a 10 percent cut, planted for anyone who misread the rate. Exact working: $847,300 \times 7/8 = 5,931,100 / 8 = 741,387.50$. Doing it as a fraction avoids the decimal multiplication altogether. Learn the fraction equivalents cold (12.5 percent is $1/8$, 16.7 percent is $1/6$, 37.5 percent is $3/8$, 62.5 percent is $5/8$) because a fraction turns most percentage items into one division.

How to practise this skill

- Memorise the fraction-to-percentage table up to eighths and twelfths. Almost every 'nasty' percentage on these tests is a clean fraction in disguise, and the fraction route is two or three times faster than decimal multiplication.
- Write the unit beside every intermediate number, including the ones you keep in your head. Most numerical errors on timed sections are unit errors, not arithmetic errors, and units are the only cheap way to catch them.
- Rehearse reverse percentages as their own drill until dividing by 1.15 feels more natural than subtracting 15 percent. It is a single, identifiable technique that accounts for a disproportionate share of lost marks.
- Before selecting, re-read the last clause of the stem out loud. 'By how much did it fall' and 'what was it after the fall' have different answers, and both are always on the option list.

- Keep an error log with four columns (misread the question, wrong operation, unit slip, arithmetic slip), and tally which one you actually commit. Three sessions is usually enough to show that one column dominates, and that is the column to train.
- Work untimed until your method is stable, then add the clock in stages: generous, then realistic, then 20 percent tighter than the real test. Attempts stay on this device, so a slow first pass costs nothing.

Glossary

Percentage point: The arithmetic difference between two percentages. A rate moving from 4 percent to 5 percent rises by one percentage point but by 25 percent in relative terms; questions specify which they want and the two answers are both on the option list.

Reverse percentage: Recovering an original amount from a figure that already includes a percentage change, by dividing by the multiplier (1.15 for a 15 percent rise, 0.88 for a 12 percent fall) rather than subtracting the percentage from the new figure.

Multiplier (decimal factor): The single number a quantity is multiplied by to apply a percentage change: 1.20 for plus 20 percent, 0.80 for minus 20 percent. Successive changes are handled by multiplying the factors, which is why plus 20 then minus 20 gives 0.96, not 1.

Ratio part: One unit of a ratio split. Dividing the total by the sum of the ratio terms gives the value of one part; every share and every difference between shares is then a whole number of parts.

Unit rate: A quantity expressed per one of something: price per litre, pages per minute, litres per 100 km. Comparisons are only valid once every option has been converted to the same unit rate.

Combined rate: The sum of two or more individual rates working simultaneously. Times cannot be added or averaged; only rates add, and the combined time is the total work divided by the combined rate.

Significant figure estimate: Rounding every input to one meaningful digit to get an approximate answer in a few seconds. Its purpose is not accuracy but elimination: it removes any option that is off by a factor or has the sign of the change wrong.

Distractor: A wrong option written deliberately to match a specific predictable error: subtracting instead of dividing, reporting the change instead of the result, averaging instead of combining rates. Recognising the pattern is often faster than recomputing.

Study this next

- Numerical reasoning skill hub: <https://learn.novusstreamsolutions.com/aptitude/skills/numerical-reasoning>
- Numerical reasoning practice in the Cognitive Skills Lab:
<https://learn.novusstreamsolutions.com/cognitive-skills/numerical-reasoning>
- Data interpretation lesson:
<https://learn.novusstreamsolutions.com/aptitude/skills/data-interpretation/lesson>
- Office numeracy suite: <https://learn.novusstreamsolutions.com/aptitude/tests/office-numeracy>
- Finance, accounting and insurance careers:
<https://learn.novusstreamsolutions.com/careers/families/finance-accounting-and-insurance>
- Numerical reasoning lesson on Novus Learn:
<https://learn.novusstreamsolutions.com/aptitude/skills/numerical-reasoning/lesson>

Where this material comes from

- All worked figures above were written and checked for Novus Learn. Every division, ratio split and rate calculation in this lesson was verified by recomputing it in the reverse direction.
- Conventions only (percentage point, unit rate, significant figures) cross-checked against standard public references such as the Wikipedia articles 'Percentage', 'Ratio' and 'Rate (mathematics)'. No item text is

drawn from any published test.

- Novus Learn aptitude construct registry (catalog seed) for construct scope and suite mapping.
- Public educational framing only: not affiliated with any official exam board, publisher or employer, and no copyrighted test item is reproduced.

Educational preparation only. Novus Learn does not administer official exams and does not guarantee scores or hiring outcomes.

Data interpretation

Tables, charts, graphs, dashboards, trends, comparisons, and evidence-based conclusions.

Data interpretation is the discipline of getting a correct number out of a table, chart or dashboard that was not built to make your question easy, and of saying so when the data cannot answer it at all. It dominates graduate and analyst screening, and it is the section where strong arithmetic still fails, because the marks are lost in the header row, the axis scale and the wording of the question. The same skill is the daily work of anyone who reports on a management pack, a clinical audit or a stock ledger. Practice is stored on this device only; there is no account and nothing is uploaded unless you export it.

What you should be able to do after this lesson:

1. Read a value correctly out of a table or chart including its units multiplier, footnotes and any 'excludes' or 'provisional' qualifier attached to the row.
2. Compute share of total, percentage change and percentage-point change from the same pair of cells, identify which of the three a question is asking for, and separate a movement in a rate from a movement in the underlying count.
3. Work with index numbers relative to a base year, including why a change of five index points is almost never a five percent change.
4. Join two tables on a shared key and produce a normalised figure (per head, per unit, per thousand) rather than comparing raw totals.
5. Recognise chart presentation effects (truncated axes, dual axes, cumulative versus periodic series), and answer from the numbers rather than from the visual impression.
6. Apply the 'cannot say' discipline: state precisely which extra fact would be needed before the question becomes answerable.

Worked examples and pitfalls

Read the header: (£000) changes every answer by a factor of a thousand: Table titled 'Regional revenue, year to March (£000)': North 1,240; South 986; East 1,455; West 719. Total = $1,240 + 986 + 1,455 + 719 = 4,400$, so the business turned over 4,400 thousand pounds, that is 4.4 million. East's share is $1,455 / 4,400 = 33.1$ percent. The whole set: North 28.2 percent, South 22.4, East 33.1, West 16.3, summing to 100. Two things go wrong here. The first is reading East's revenue as 1,455 pounds and then reporting a business with a total turnover of 4,400 pounds, which nobody notices because every option is scaled the same way, until the question asks for revenue in millions and only one option is right. The second is the comparison wording. East's share is $(1,455 - 719) / 4,400 = 16.7$ percentage points above West's, and East's revenue is $(1,455 - 719) / 719 = 102$ percent more than West's, that is slightly more than double. 'Sixteen point seven' and 'a hundred and two' both describe the same two cells honestly, and the question decides which one is correct. Note that the percentage-point figure must be computed from the unrounded shares rather than by subtracting the rounded ones, or the last digit will not survive.

One pair of rows, three correct increases: A complaints table: 2023, 120,000 orders, complaint rate 4.0 percent; 2024, 150,000 orders, complaint rate 5.0 percent. Three defensible answers to 'how much did complaints increase?'. The rate rose by 1.0 percentage point. The rate rose by $(5.0 - 4.0) / 4.0 = 25$ percent

in relative terms. And the count of complaints rose from $0.04 \times 120,000 = 4,800$ to $0.05 \times 150,000 = 7,500$, which is $(7,500 - 4,800) / 4,800 = 56.25$ percent. All three are arithmetically right; only one answers the question in front of you. The pattern to internalise is that a rate and a count move together only when the denominator is fixed, and here it is not. Order volume grew 25 percent as well. If the question is about customer experience, the rate is the honest figure; if it is about how many complaint handlers to hire, the count is. Test items usually ask for the one you would not have chosen.

Index numbers: five points is not five percent: A cost index with 2020 = 100 reads 104 in 2021, 111 in 2022 and 109 in 2023. From 2021 to 2023 the index rose 5 points, but the percentage change is $5 / 104 = 4.8$ percent, because the base for the comparison is 104, not 100. From 2022 to 2023 it fell 2 points, which is $-2 / 111 = -1.8$ percent, and note that costs fell even though the index remains 9 percent above the 2020 base. A level and a change are different claims. The only comparison where points and percent coincide is against the base year itself: 2020 to 2023 is 100 to 109, exactly plus 9 percent. Watch also for a rebased series, where a table switches to 2022 = 100 partway down; the two segments cannot be compared directly without converting one of them, and an item that quietly rebases is testing whether you read the column heading.

Joining two tables: totals and per-head figures disagree on purpose: Table 1, headcount by site: Leeds 84, Derby 47, Bristol 129. Table 2, absence days recorded in the same period: Leeds 630, Derby 300, Bristol 903. 'Which site has the worst absence problem?' On raw totals Bristol is worst at 903 days. Normalise per head and the ranking changes: Leeds $630 / 84 = 7.50$ days per employee, Bristol $903 / 129 = 7.00$, Derby $300 / 47 = 6.38$. Leeds is worst, Bristol is merely biggest. The organisation-wide figure is $1,833 / 260 = 7.05$ days per head, which is a useful reference line: Leeds is above it, the other two below. The general rule is that any comparison between units of different size demands a denominator, and the denominator has to come from the other table. Items are built so the raw-total answer and the per-head answer are both on the option list, and so the site with the biggest total is never the site with the highest rate.

The truncated axis: measure the numbers, not the bars: A quarterly satisfaction chart with a y-axis running from 78 to 82 shows bars at 79.2, 79.8, 80.4 and 81.1. Visually the last bar looks several times taller than the first, because only the top 4 points of a 100-point scale are drawn. The actual movement is $81.1 - 79.2 = 1.9$ points, which on the score's own scale is a relative rise of $1.9 / 79.2 = 2.4$ percent. If the question asks 'by approximately what percentage did satisfaction improve', the answer is about 2 percent, and the distractor built from the bar heights will be something like 40 or 400 percent. Related presentation effects to check before answering: a dual-axis chart where two series use different scales and appear to cross meaningfully when they do not; a cumulative series, where a flattening line still means the total is growing, just more slowly; and a logarithmic axis, where equal vertical distances are equal ratios rather than equal amounts. In every case the defence is the same. Find the printed numbers, or read the gridline values, and compute.

Cannot say: revenue is not profit: A product table shows units sold and total revenue. Product P: 4,200 units, 71,400 pounds. Product Q: 1,800 units, 41,400 pounds. Average selling price is $71,400 / 4,200 = 17.00$ for P and $41,400 / 1,800 = 23.00$ for Q, so Q earns more per unit while P earns more in total. Now the statement to evaluate: 'P is more profitable than Q.' The correct response is cannot say. Profit needs cost, and the table has no cost column; a product with a 17 pound price and a 16 pound unit cost is less profitable than one priced at 23 with a cost of 9, and nothing here rules that out. Contrast with 'Q generated more revenue per unit than P', which the table fully supports and which is true. The habit worth building is to finish every cannot-say judgement with the missing input named out loud ('cannot say, because unit cost is not given') because that forces you to distinguish a genuinely unanswerable item from one you simply have not worked hard enough on.

How to practise this skill

- Read the title, the units line, the row and column headers and any footnote before you look at a single value. Roughly the first fifteen seconds of an item should contain no arithmetic at all, and that fifteen seconds is what prevents the thousand-fold and percentage-point errors.

- For every item, write down which of the three quantities is wanted (share of total, relative change, or change in percentage points), before computing. Most wrong answers on this construct are correct arithmetic applied to the wrong quantity.
- Whenever two groups differ in size, ask what the denominator should be. If a question compares sites, teams, countries or periods of unequal length and you have not divided by something, you are almost certainly answering the wrong question.
- Practise the cannot-say items separately and force yourself to name the missing variable each time. Candidates who train only on computational items reliably over-answer inference statements under time pressure.
- Do not redraw or re-scale charts in your head. Locate the gridline values or the data labels, and if neither exists, interpolate between two labelled gridlines and state the bound rather than guessing a point value.
- Time yourself per item rather than per section. Data interpretation sets share a stimulus, so the first item costs the reading time and the rest should be fast; if item four takes as long as item one, you did not build a mental map of the table.

Glossary

Units multiplier: A scaling note in a table title or column header, such as (£000), (millions) or (per 1,000 population). It applies to every value in scope and is the single most common source of order-of-magnitude errors.

Index number: A series rescaled so a chosen base period equals 100. Changes between two non-base periods must be divided by the earlier value, so a movement in index points is not a percentage change except when measured from the base.

Rebasing: Restating an index against a new base period. Segments of a series with different bases cannot be compared directly, and a table that rebases partway down is testing whether you read the headings.

Truncated axis: A chart whose value axis does not start at zero, which exaggerates the apparent size of differences between bars or points. Legitimate for showing small movements in a large quantity, misleading if read as area or height.

Cumulative series: A line showing a running total rather than each period's value. It can only go up or stay flat, so a flattening cumulative line means the periodic figure is falling, not that the total is.

Weighted average: An average in which each value is multiplied by the size of the group it represents.

Averaging two group percentages directly is only correct when the groups are the same size, which in these tables they rarely are.

Normalisation: Dividing a raw figure by an exposure measure (headcount, units sold, population, days open), so groups of different size can be compared. The denominator usually lives in a second table.

Cannot say: The verdict when a statement is neither supported nor contradicted by the data supplied. A correct cannot-say answer can always be defended by naming the specific missing variable.

Study this next

- Data interpretation skill hub: <https://learn.novusstreamsolutions.com/aptitude/skills/data-interpretation>
- Data interpretation practice in the Cognitive Skills Lab:
<https://learn.novusstreamsolutions.com/cognitive-skills/data-interpretation>
- Numerical reasoning lesson:
<https://learn.novusstreamsolutions.com/aptitude/skills/numerical-reasoning/lesson>
- Financial data interpretation suite:
<https://learn.novusstreamsolutions.com/aptitude/tests/financial-data-interpretation>
- Scientific data interpretation suite:
<https://learn.novusstreamsolutions.com/aptitude/tests/scientific-data-interpretation>

- Data interpretation lesson on Novus Learn:
<https://learn.novusstreamsolutions.com/aptitude/skills/data-interpretation/lesson>

Where this material comes from

- Every table, index series and chart described above was constructed for Novus Learn, and each figure was verified by recomputing the totals and the reverse calculation.
- Definitions of index numbers, rebasing and weighted averages cross-checked against standard public references such as the Wikipedia articles 'Index (economics)' and 'Weighted arithmetic mean'. Terminology only; no data or item text is taken from any source.
- Novus Learn aptitude construct registry (catalog seed) for construct scope and suite mapping.
- Public educational framing only: not affiliated with any official exam board, publisher or employer, and no copyrighted test item is reproduced.

Educational preparation only. Novus Learn does not administer official exams and does not guarantee scores or hiring outcomes.

Critical thinking

Evaluating evidence, assumptions, arguments, credibility, and alternative explanations.

Critical thinking assessments ask you to take an argument apart: what is being concluded, what it rests on, what evidence would settle it, and whether a proposed conclusion actually follows. Published critical-thinking batteries typically run five task types - inference, recognition of assumptions, deduction, interpretation, and evaluation of arguments - and they are deliberately built so that agreeing with a conclusion and judging the argument as strong come apart. The construct is heavily weighted in policy analysis, audit, legal, investigative and graduate selection, and it is the same skill that stops a plausible chart from turning into a bad decision. Practice stays on this device; there is no account and nothing is uploaded unless you export it.

What you should be able to do after this lesson:

1. State an argument's conclusion in your own words before evaluating it, and identify which sentences are premises and which are background.
2. Use the negation test to separate a required assumption from a statement that would merely strengthen the argument.
3. Distinguish validity from truth, and identify affirming the consequent as distinct from the valid modus tollens form.
4. Compute a posterior probability on a screening example and explain why the rate of true positives among all positives is far lower than intuition suggests.
5. Generate at least two alternative explanations - selection, reverse causation, a common third factor - for any claimed causal effect, and name the comparison that would rule them out.
6. Judge argument strength on relevance and directness rather than on agreement, including marking arguments you personally reject as strong.

Worked examples and pitfalls

The negation test finds assumptions; nothing else does: Argument: 'The council should fit bin sensors across the district. In the trial depot they cut collection trips by a fifth.' Candidate assumption A: 'The trial depot's waste pattern is broadly representative of the rest of the district.' Negate it - suppose the trial depot is nothing like the rest of the district - and the argument collapses, because a result that does not transfer supports nothing about the district. So A is a required assumption. Candidate assumption B: 'Bin sensors are the cheapest available technology for this purpose.' Negate it - suppose they are the most expensive - and the argument still stands as given, since it argued from a reduction in trips, not from cost. B is not required; it

would strengthen a cost-based argument that was never made. Candidate C: 'Fewer collection trips reduce total operating cost.' Negate it and the recommendation loses its point, so C is required too, even though the argument never mentions cost. That last case is the one candidates miss, because the required assumption is what bridges the evidence to the recommendation, and bridging assumptions are by definition unstated. The failure mode across this whole item type is picking statements that support the conclusion instead of statements the argument cannot survive without.

Base rates: 90 percent accurate, 15 percent right: A production line runs an automated defect test. Defects occur in 1 percent of units. The test flags 90 percent of genuinely defective units, and wrongly flags 5 percent of good units. A unit has just been flagged - how likely is it to be defective? Work in whole units rather than probabilities. Take 10,000 units: 100 are defective and 9,900 are good. Of the 100 defective, the test flags 90. Of the 9,900 good, it wrongly flags 5 percent, which is 495. Total flags: $90 + 495 = 585$. Of those, 90 are genuine, so the answer is $90 / 585 = 15.4$ percent. Better than the 1 percent base rate, and nowhere near the 90 percent that intuition offers. The error has a name - confusing the probability of a flag given a defect with the probability of a defect given a flag - and it has a practical consequence: the rework queue has to be sized for 585 units a batch, not 100, and 495 of the 585 units sitting in it are perfectly good. Whenever a test, screening rule or model accuracy figure appears in an item, ask how many of the negatives there are, because with a rare condition the false positives from a large clean population swamp the true positives from a small affected one.

Three alternative explanations, and the comparison that settles it: Claim: 'Employees who attended the optional resilience workshop took 30 percent fewer sick days last year, so the workshop works.' Alternative one, selection: attendance was optional, so the people who signed up may have been the healthier and more engaged to begin with, and would have taken fewer sick days regardless. Alternative two, reverse direction: employees who were frequently unwell were the least able to attend a full-day workshop, so illness determined attendance rather than the other way round. Alternative three, a common third factor: the workshop ran on Thursday afternoons at head office, so attendance largely tracks being an office-based rather than shift-based worker, and shift workers take more sick leave for reasons that have nothing to do with resilience training. What would settle it: randomly assigning the offer to half the workforce and comparing the two groups. Failing that, the cheapest useful check already exists in the payroll data - compare attendees' and non-attendees' sick days in the year before the workshop. If attendees were already 30 percent lower, selection is doing the entire job. The habit worth building is to name the comparison group before you accept any before-and-after number.

Deduction: affirming the consequent versus modus tollens: Rule: 'If an invoice is over 10,000 euro, it requires two signatures.' Argument A: 'This invoice has two signatures, therefore it is over 10,000 euro.' Invalid - this is affirming the consequent. The rule makes two signatures necessary for large invoices; it says nothing that prevents a cautious manager from double-signing a 400 euro invoice, so a two-signature invoice of any size is consistent with the rule. Argument B: 'This invoice has only one signature, therefore it is not over 10,000 euro.' Valid - this is modus tollens, denying the consequent to deny the antecedent, and it holds assuming the rule was followed. Argument C: 'This invoice is not over 10,000 euro, therefore it does not have two signatures.' Invalid again, denying the antecedent. Two further points that deduction items test directly. First, validity is about form alone: 'All banks close on Sundays; Riverton Mutual is a bank; therefore Riverton Mutual closes on Sundays' is perfectly valid even if the first premise is false, and a valid argument with a false premise can deliver a false conclusion. Second, in these items you must accept the premises as given even when you know them to be untrue, because the question is whether the conclusion follows, not whether it is true.

Interpretation: what a survey figure does and does not license: Data: all 1,200 employees of a firm were surveyed. 62 percent said they would accept a four-day week at 90 percent pay. Of those who said yes, 71 percent were under 35. Conclusion one: 'A majority of this firm's employees would accept the trade.' This follows - 62 percent of a complete census of the firm is a majority, and the conclusion is properly limited to

this firm. Conclusion two: 'A majority of the firm's under-35 employees would accept the trade.' This does not follow, and the arithmetic shows why. Sixty-two percent of 1,200 is 744 accepters, and 71 percent of 744 is about 528 accepters aged under 35. That is the composition of the accepters, not the acceptance rate within the under-35 group, and without knowing how many under-35s the firm employs the rate is undetermined: if there are 900, the rate is $528 / 900 = 59$ percent; if there are 600, it is 88 percent; if there were 1,100 the rate would be below half. Conclusion three: 'The firm should move to a four-day week.' This does not follow either - a stated preference about pay says nothing about output, shift cover, customer hours or cost. Notice that conclusion two is the base-rate error from the screening example, wearing survey clothing.

Strong and weak arguments on the same question: Question: should the agency publish raw inspection scores for every site? Strong argument for: 'In the neighbouring authority, publication was followed by a measurable fall in repeat violations, which is the stated aim of the inspection regime.' It is directly about this decision, it addresses the regime's own objective, and it offers evidence - and if true it would change your view. Weak argument for: 'Yes, because transparency is always good.' Sweeping, unsupported, and it would apply identically to publishing anything at all, which is the tell. Strong argument against: 'Raw scores omit the severity weighting, so a site with one critical failure can rank above a site with six minor ones, and a member of the public cannot see the difference.' Specific, mechanistic, directly about the proposal, and it would change your view. Weak argument against: 'No, because businesses will object.' It may be true and it may matter politically, but it does not bear on whether publication achieves the regime's aim, and no evidence is offered. Two tests separate strong from weak: is the argument about the exact question asked, and would accepting it change the decision? Note that you can hold a firm personal view on publication and still be required to mark the argument on the other side as the strong one - that split is what the item type is measuring.

How to practise this skill

- Negate every candidate assumption out loud. If the argument survives the negation, it was never an assumption, however supportive it sounds. This single habit converts the assumption item type from guesswork into a mechanical check.
- Write the conclusion in your own words before reading any option. Half of the wrong answers on inference and evaluation items are responses to a conclusion the argument never reached.
- Ask 'percentage of what?' on every percentage in the stimulus and reconstruct the denominator in whole units. Base-rate and composition errors both dissolve the moment you write out counts instead of rates.
- Deliberately practise marking arguments you disagree with as strong and arguments you agree with as weak. Assessments in this construct are built to catch agreement masquerading as evaluation, and the effect is largest on politically loaded stimuli.
- For any causal claim, write two alternatives - selection and a third factor - and name the comparison group that would rule them out, before you decide whether the evidence supports the claim.
- Log wrong answers by task type rather than by topic: inference, assumption, deduction, interpretation, evaluation. Candidates are rarely weak across all five, and the profile tells you where the next hour of practice belongs.

Glossary

Conclusion: The claim an argument is trying to establish. It is not always last, is often signalled by therefore, so, or should, and locating it correctly determines every subsequent judgement about the argument.

Assumption: An unstated premise the argument requires in order to work. Identified by negation: negate it and a genuine assumption brings the argument down, while a merely helpful statement leaves it standing.

Validity and soundness: An argument is valid when the conclusion cannot be false while the premises are true, which is a property of form alone. It is sound when it is valid and the premises are actually true.

Affirming the consequent: The invalid pattern 'if P then Q; Q; therefore P'. It is the most common deductive

distractor because it differs from the valid modus tollens form by only the position of a negation.

Base rate: How common something is in the population before any test or evidence is applied. Ignoring it makes accurate-sounding tests appear far more informative than they are when the condition is rare.

Confounder: A third factor associated with both the supposed cause and the outcome, capable of producing the entire observed relationship on its own. Ruled out by randomisation or by an explicit comparison group.

Selection effect: A difference between groups created by how people entered them rather than by the treatment under study. Optional programmes and voluntary surveys are where it appears most often.

Falsifiability: The property of a claim that some observation could show it to be wrong. A claim compatible with every possible result carries no information, which is why 'what would change your mind?' is a diagnostic question.

Study this next

- Critical thinking skill hub: <https://learn.novusstreamsolutions.com/aptitude/skills/critical-thinking>
- Practise the critical thinking bank: <https://learn.novusstreamsolutions.com/cognitive-skills/critical-thinking>
- Verbal reasoning lesson: <https://learn.novusstreamsolutions.com/aptitude/skills/verbal-reasoning/lesson>
- Government audit suite: <https://learn.novusstreamsolutions.com/aptitude/tests/government-audit>
- Investigator and detective reasoning suite:
<https://learn.novusstreamsolutions.com/aptitude/tests/investigator-and-detective-reasoning>
- Critical thinking lesson on Novus Learn:
<https://learn.novusstreamsolutions.com/aptitude/skills/critical-thinking/lesson>

Where this material comes from

- Worked items written for Novus Learn. The arguments, survey figures, defect-test numbers and evaluation options above are original and invented for this lesson; no published or copyrighted test item is reproduced.
- The five task types named in the summary - inference, assumption, deduction, interpretation, evaluation of arguments - describe a structure used publicly across several critical-thinking assessments; no affiliation with any publisher is claimed or implied.
- Terminology follows standard, widely published usage in introductory logic and research methods - validity, soundness, affirming the consequent, base rate, confounder, selection effect.
- Novus Learn aptitude construct registry (catalog seed) for the construct scope and the suite mapping shown in the related links.
- Public educational framing only - not affiliated with any official exam board, publisher or employer.

Educational preparation only. Novus Learn does not administer official exams and does not guarantee scores or hiring outcomes.

Verbal reasoning

Drawing conclusions from written information without relying on outside assumptions.

Verbal reasoning is the discipline of answering strictly from a passage: deciding whether a statement is True, False, or Cannot Say on the evidence given, and refusing to import anything you happen to know about the topic. It is the workhorse construct in graduate and civil-service screening, banking and consulting sifts, and language-focused public-service batteries, where a passage about parcel volumes or grant conditions is followed by four statements to classify. The same discipline is what stops a contract review, a grant assessment or an incident summary from quietly acquiring facts nobody wrote down. Practice here is device-local: no account, and nothing leaves this device unless you export it.

What you should be able to do after this lesson:

1. Classify a statement as True, False, or Cannot Say, and state in one sentence which words in the passage force the classification.
2. Explain the difference that costs the most marks: 'Cannot Say' means undetermined by the passage, while 'False' means contradicted by it.
3. Spot quantifier drift between passage and statement (all, some, most, only, none) and show why 'all A are B' never licenses 'all B are A'.
4. Refuse a causal upgrade: recognise when a passage reports co-occurrence or sequence and a statement claims cause.
5. Separate what a passage asserts from what it reports somebody else asserting, and classify statements about each correctly.
6. Do the small arithmetic hidden inside verbal items - percentage changes, fractions of a stated total - without leaving the passage for outside data.

Worked examples and pitfalls

The three-way rule on one short passage: Passage: 'The Riverton depot handled 4,200 parcels in March, a 12 percent increase on February. A night shift was introduced in January; in March it processed 40 percent of the parcels the depot handled.' Now take four statements. (1) 'The depot handled 3,750 parcels in February.' February is March divided by 1.12: $4,200 / 1.12 = 3,750$ exactly, so this is True - the passage determines it even though the number is not printed. (2) 'The night shift caused the March increase.' Cannot Say. The passage puts the night shift and the increase in the same depot in the same period; it never claims one produced the other, and it never denies it either. Most candidates mark this False, reasoning correctly that correlation is not causation but then choosing the wrong label: False is reserved for statements the passage contradicts. (3) 'The night shift processed more than 1,700 parcels in March.' Forty percent of 4,200 is 1,680, which is not more than 1,700, so this is False - contradicted by arithmetic on the passage's own figures. (4) 'The depot handled more parcels in March than in January.' Cannot Say: January is mentioned only as the month the shift started, and no January volume is given.

Quantifier drift and illicit conversion: Passage: 'All accredited suppliers submit quarterly audits. Some suppliers in the northern region are accredited.' Statement A: 'Some suppliers in the northern region submit quarterly audits.' True, and it is worth seeing why the chain holds: at least one northern supplier is accredited, and every accredited supplier submits, so at least one northern supplier submits. Statement B: 'Every supplier that submits quarterly audits is accredited.' This reverses the first sentence. 'All accredited suppliers submit' leaves the door wide open for unaccredited suppliers to submit as well - perhaps voluntarily, perhaps under another rule - so the answer is Cannot Say. Turning 'all A are B' into 'all B are A' is called illicit conversion and it is the single most productive trap in this format, because the reversed sentence sounds like a paraphrase. Statement C: 'No unaccredited supplier submits quarterly audits' is the same reversal wearing a negative coat, and it is Cannot Say for the same reason. Statement D: 'All northern suppliers submit quarterly audits' is also Cannot Say: 'some are accredited' says nothing about the rest.

Asserted versus reported: Passage: 'The committee's report claims that the scheme cut waiting times by a third. Two of the four regional boards have disputed that figure.' Statement A: 'The scheme cut waiting times by a third.' Cannot Say. What the passage asserts is that a report claims this; the claim's truth is never settled, and the dispute does not settle it either. Statement B: 'At least two regional boards disagree with the report's waiting-time figure.' True, and note 'at least' - the passage says two disputed it, and two of four disputing is consistent with 'at least two'. Statement C: 'All four regional boards accept the figure.' False, directly contradicted. Statement D: 'The report is wrong.' Cannot Say. This item type appears constantly in policy and diplomacy passages, where a paragraph stacks a claim, a source, and a reaction, and the reader has to keep three different truth-conditions apart. The reliable habit is to underline the reporting verb - claims, argues, estimates, alleges, projects - and remember that everything downstream of it belongs to the source, not to the passage.

Percentages are not symmetric: Passage: 'Membership fell from 8,000 to 6,000 over two years, then recovered by 25 percent in the third year.' Statement: 'Membership at the end of year three was higher than at the start.' The fall is 2,000 on a base of 8,000, which is 25 percent. The recovery is 25 percent of the new base: $6,000 \times 1.25 = 7,500$. That is below 8,000, so the statement is False. The trap is elegant, because a 25 percent fall followed by a 25 percent rise feels like it should cancel. It does not: to get back from 6,000 to 8,000 you need a 33.3 percent rise, since $2,000 / 6,000 = 0.333$. Whenever a verbal item states a change as a percentage, write down what the percentage is a percentage of before you touch it. A related version supplies the recovery as an absolute number instead - 'recovered by 2,000 members' - which does return the total to 8,000, and candidates who answered the first version from memory get the second one wrong.

Verbal analogies: name the relation as a sentence: Item: 'ANTISEPTIC is to INFECTION as ____ is to ____.' Options: (a) vaccine : immunity, (b) insulation : heat loss, (c) medicine : illness, (d) bandage : wound. Write the stem relation as a full sentence first: 'An antiseptic is applied in order to prevent an infection from occurring.' Now test each option against that exact sentence. Option (a) runs the other way - a vaccine is given to produce immunity, not to prevent it. Option (c) is the right family but the wrong specificity: medicine typically treats an illness that has already started. Option (d) is applied after the wound exists. Option (b) fits precisely: insulation is applied in order to prevent heat loss from occurring. The general method is the same on the simpler item types. 'BOOK is to LIBRARY as PAINTING is to ____' has canvas (the material), artist (the maker), and gallery (the place a collection is kept and shown) among its options; all three are genuine relations to 'painting', and only the third matches the stem sentence. Options in analogy items are chosen to be true statements about the words, just not the stated relation.

Two premises with 'some' prove nothing: Argument: 'Some depot managers are qualified engineers. Some qualified engineers hold a rigging certificate. Therefore some depot managers hold a rigging certificate.' This is invalid, and the way to prove it is a counter-model rather than an argument. Let the depot managers be Ana and Ben; let the qualified engineers be Ben and Cara; let the certificate holders be Cara alone. Premise one holds: Ben is a manager and an engineer. Premise two holds: Cara is an engineer with the certificate. The conclusion fails: neither Ana nor Ben holds the certificate. Since a single consistent world makes both premises true and the conclusion false, the argument cannot be valid, so in a True / False / Cannot Say framing the conclusion is Cannot Say. Two premises that both begin 'some' never combine into a conclusion, because 'some' gives you an overlap without telling you where the overlap sits. Contrast a valid pair: 'All accredited labs are inspected annually. No inspected facility is exempt from the fee.' Every accredited lab is inspected, and no inspected facility is exempt, so no accredited lab is exempt - True.

How to practise this skill

- Answer from the passage even when you know the subject. Candidates with a background in the topic score worse on verbal items than they expect, because their own knowledge quietly supplies the missing premise that turns a Cannot Say into a True.
- Run the two-worlds test on every candidate 'Cannot Say': can you imagine a world where all the passage's sentences hold and the statement is true, and another where they hold and it is false? If both, it is Cannot Say. If only the false world exists, it is False.
- Underline the quantifiers and hedges in the statement (all, some, only, most, may, must, likely) and find their counterparts in the passage. Roughly half of the wrong answers in this construct come from a single word swapped between the two.
- Budget about 25 to 30 seconds per statement rather than per passage, and read the passage once for structure before touching the statements. Re-reading the whole passage for each statement is what causes candidates to run out of time on the last set.
- Keep a wrong-answer log with one label per error: quantifier, causation, reported-versus-asserted, arithmetic, outside knowledge. A clear pattern almost always emerges within about forty items, and it is usually a single label.
- Work untimed until the three-way rule is automatic, then add the clock. Attempts are stored on this device

only, so a slow first pass through a passage set costs you nothing but the time you spend on it.

Glossary

Entailment: A statement is entailed by a passage when it cannot be false while every sentence of the passage is true. Entailment is the only thing that earns a 'True' in this format; being plausible, likely, or well known does not.

Cannot Say: The verdict for a statement the passage neither entails nor contradicts. It is a claim about the passage, not about the world, which is why a statement you know to be true in real life can still be Cannot Say.

Quantifier: A word fixing how much of a group a claim covers: all, most, some, few, none, only. Swapping one quantifier for another changes the logical content completely while barely changing how the sentence reads.

Illicit conversion: The invalid move from 'all A are B' to 'all B are A', or from 'if P then Q' to 'if Q then P'. It preserves the words and destroys the logic, which is why converted sentences make such effective wrong answers.

Counter-model: A concrete, consistent scenario in which the premises hold and the conclusion fails. Producing one is the fastest possible proof that an argument is invalid, and it takes two or three named individuals.

Hedge: A qualifier such as may, could, is expected to, or is associated with. A hedged sentence in a passage cannot support an unhedged statement, and an unhedged passage sentence is not weakened by a hedged statement.

Reporting verb: A verb such as claims, argues, estimates, or alleges that attributes the following content to a source. Everything downstream of it is the source's assertion, and the passage takes no position on it.

Analogy relation: The specific link between the two stem words in an analogy item - tool and user, item and container, action and purpose. Naming it as a full sentence before reading the options removes almost all of the guesswork.

Study this next

- Verbal reasoning skill hub: <https://learn.novusstreamsolutions.com/aptitude/skills/verbal-reasoning>
- Practise the verbal reasoning bank:
<https://learn.novusstreamsolutions.com/cognitive-skills/verbal-reasoning>
- Critical thinking lesson: <https://learn.novusstreamsolutions.com/aptitude/skills/critical-thinking/lesson>
- Reading comprehension lesson:
<https://learn.novusstreamsolutions.com/aptitude/skills/reading-comprehension/lesson>
- General civil-service aptitude suite:
<https://learn.novusstreamsolutions.com/aptitude/tests/general-civil-service-aptitude>
- Verbal reasoning lesson on Novus Learn:
<https://learn.novusstreamsolutions.com/aptitude/skills/verbal-reasoning/lesson>

Where this material comes from

- Worked items written for Novus Learn. Every passage, statement set and figure above is original and invented for this lesson; no published or copyrighted test item is reproduced.
- Terminology follows standard, widely published usage in introductory logic and assessment writing - entailment, quantifier, illicit conversion, counter-model.
- Novus Learn aptitude construct registry (catalog seed) for the construct scope and the suite mapping shown in the related links.

- Public educational framing only - not affiliated with any official exam board, publisher or employer.

Educational preparation only. Novus Learn does not administer official exams and does not guarantee scores or hiring outcomes.

Problem solving

Selecting methods, combining information, troubleshooting, and reaching practical solutions.

Problem solving is the construct that asks you to choose a method, not just execute one: to combine a table with a rule, decide whether an estimate settles the question or an exact calculation is needed, narrow a fault by halving the search space, and check the answer against the constraint that actually binds. It shows up across operations, logistics, manufacturing, technician and analyst selection, and in the situational sections of public-safety batteries. It is also the construct where the marking rewards a defensible route as much as a number. Everything you practise stays on this device unless you export it.

What you should be able to do after this lesson:

1. Combine a pricing or specification table with a written brief to reach an answer neither source contains on its own, and find the break-even point where the better option changes.
2. Narrow a fault on a chain of components by bisection, and state the worst-case number of tests before you begin.
3. Solve a working-backwards problem from a known end state, then verify by running the sequence forwards.
4. Decide whether an order-of-magnitude estimate settles a feasibility question or whether the margin is close enough to require exact arithmetic.
5. Identify the binding constraint in a resource-allocation problem and explain why a per-unit heuristic on the non-binding resource gives the wrong answer.
6. Separate a root cause from a symptom, and describe the test that would confirm the cause rather than merely correlate with it.

Worked examples and pitfalls

Two price structures and the distance where they cross: A delivery of 55 km. Courier A charges 8.00 base plus 0.45 per km. Courier B charges 20.00 flat for the first 40 km, then 0.90 per km beyond that. Which is cheaper? Courier A costs 8 plus 0.45 times 55, which is 8 plus 24.75, giving 32.75. Courier B costs 20 plus 0.90 times 15, which is 20 plus 13.50, giving 33.50. Courier A wins by 0.75. Now the question the item is really testing: is A always cheaper? At 45 km, A costs 8 plus 20.25 which is 28.25, while B costs 20 plus 4.50 which is 24.50. B wins comfortably. So the answer flips somewhere between, and finding where is one line of algebra. Above 40 km, B costs 20 plus 0.9 times distance minus 40, which simplifies to $0.9d$ minus 16. Set that equal to A's 8 plus $0.45d$: $8 + 0.45d = 0.9d - 16$, so 24 equals $0.45d$, so d is about 53.3 km. Below 53.3 km B is cheaper; above it A is. Check at the crossover: A is 8 plus 24 which is 32.00, and B is 48 minus 16 which is also 32.00. The general lesson is that whichever option has the lower per-unit rate always wins eventually, regardless of the base charges, and the base charges only decide where eventually starts. A single quoted distance never answers a which-is-cheaper question for a fleet.

Halving the search space beats walking the chain: Forty sensors sit on a single daisy-chained cable and exactly one connection is broken, cutting off everything downstream. How many tests do you need in the worst case? Walking the chain from one end and testing each sensor in turn takes up to 40 tests and 20 on average. Test the midpoint instead. If sensor 20 responds, the break is in the upper half and you have eliminated 20 candidates in one measurement; if it does not, the break is below and you have eliminated the other 20. Repeat on the surviving half: 40 becomes 20, then 10, then 5, then 3, then 2, then 1. Six tests, worst case, because each test halves what is left and two to the power six is 64, comfortably more than 40.

The saving grows as the problem grows: 1,000 candidates need only ten tests. Two conditions have to hold for bisection to be valid, and stating them is part of the answer. The fault must be monotone, meaning everything on one side of the break behaves differently from everything on the other; and a single test at any point must tell you which side you are on. When there are two independent breaks, or when a test is only meaningful at the ends, bisection does not apply and a different strategy is needed. Candidates who reach for bisection reflexively on a problem that fails those conditions lose more than they save.

Working backwards from the number you were given: A team started a quarter with an unknown budget. In week one it spent half of it. In week two it spent 400 of what remained. In week three it spent a third of what remained after that. It finished with 1,200. How much did it start with? Attempting this forwards means carrying an unknown through three operations. Backwards it is arithmetic. After week three, 1,200 is what is left having spent a third, so 1,200 represents two thirds of the week-three opening balance, which was therefore 1,800. Before week two's spend of 400, the balance was 1,800 plus 400, which is 2,200. That 2,200 is what remained after spending half, so the starting budget was 4,400. Verify forwards, always: 4,400 less half is 2,200; less 400 is 1,800; less a third of 1,800, which is 600, leaves 1,200. Correct. The mechanical rule is to invert each operation and apply the inversions in reverse order. The inverse of spending a third is dividing by two thirds, not multiplying by three, and that specific slip is the most common wrong answer in this family. Working backwards is the right tool whenever the end state is known exactly and the operations are individually invertible, which covers most budget, mixture and journey problems phrased as how much did it start with.

When an estimate is enough, and when it is not: A maintenance window is three hours. Two technicians must service 1,850 units, each taking about 4.5 minutes, plus a shared 20-minute setup and 15-minute teardown. Feasible? Estimate first: 1,850 units at 4.5 minutes is 8,325 technician-minutes; split between two people that is about 4,163 minutes, which is roughly 69 hours. The window is 3 hours. The answer is no by a factor of more than twenty, and no refinement of the setup and teardown figures could change it. Precision here would be wasted effort. Now the same problem with 90 units. Ninety at 4.5 minutes is 405 technician-minutes, halved to 202.5 minutes, plus 35 minutes of setup and teardown, giving 237.5 minutes. That is 3 hours and 57 and a half minutes against a 3-hour window. Still no, but only just, and now every assumption matters: whether the technicians can genuinely work in parallel, whether setup is shared or duplicated, whether 4.5 minutes is a mean or a best case. The judgement being assessed is knowing which regime you are in. If your rough figure misses the target by an order of magnitude, stop and answer. If it lands within a factor of two, the estimate has not settled anything and you must do the exact arithmetic and name your assumptions.

The binding constraint decides, and it is rarely the obvious one: A van carries at most 900 kg and at most 6 cubic metres. Pallet type X weighs 150 kg, occupies 0.8 cubic metres and is worth 200. Pallet type Y weighs 90 kg, occupies 1.4 cubic metres and is worth 260. What mix maximises value? The instinctive heuristic is value per kilogram: X gives 200 over 150, about 1.33, while Y gives 260 over 90, about 2.89, so load Y. That is wrong, and the reason is that weight is not what runs out. Value per cubic metre tells the opposite story: X gives 250 and Y gives about 186. Enumerate the feasible whole-pallet mixes. Six X uses 900 kg and 4.8 cubic metres, worth 1,200. Five X and one Y uses 840 kg and 5.4 cubic metres, worth 1,260. Four X and two Y uses 780 kg and exactly 6.0 cubic metres, worth 1,320. Three X and three Y needs 6.6 cubic metres and does not fit. Two X and three Y uses 570 kg and 5.8 cubic metres, worth 1,180. Four Y alone uses 5.6 cubic metres and is worth 1,040. The best mix is four X and two Y at 1,320. Look at what binds it: volume is used to the last cubic metre while 120 kg of payload goes unused. Once volume is identified as the binding constraint, X's superior value per cubic metre explains the X-heavy answer, and the spare weight explains why two Y still get on board. The transferable move is to compute the usage of every constraint at your proposed answer and see which one is exhausted; a per-unit ratio computed against a non-binding resource is worse than no heuristic at all.

Cause, symptom, and the test that tells them apart: A pump trips out twice a shift. The obvious fix is to reset

it, which works for about four hours. Ask why once and you find the thermal overload relay is tripping. Ask again and the motor is running hot. Again and the bearing is running hot. Again and the grease has dried out. Again and the lubrication interval was set for a duty cycle far lighter than the pump has actually been running since the line was rebalanced last year. Each level supports a different fix: resetting the trip costs nothing and lasts four hours; replacing the relay costs a part and lasts until the bearing seizes; regreasing lasts weeks; changing the lubrication schedule to match the real duty cycle is the one that stops the fault recurring. Stop asking why when the next answer stops being something you can act on, or when it leaves the boundary of the system you can change. There is one discipline that separates this from storytelling. A correlation is a lead, not a cause: the fact that the trips started after the line was rebalanced is suggestive, not proof. The confirming test is to intervene and observe: restore the original duty cycle, or apply the corrected greasing interval to this pump and not to the identical one on the parallel line, and see which one trips. If a proposed cause cannot be tested by changing it, treat it as a hypothesis and say so rather than closing the investigation.

How to practise this skill

- Write the question you are actually being asked at the top of your working, in your own words. Problem-solving items usually supply more information than the question needs, and the commonest failure is a correct calculation of something nobody asked for.
- Before any comparison of two options, ask where they cross. One quoted quantity never settles which is cheaper or faster, and the break-even is usually a single line of algebra away.
- State the worst-case number of steps before starting a search or a fault trace. If bisection applies, say so and say why; if the fault is not monotone or a test only works at the ends, say that instead and pick a different method.
- Estimate before you compute, then decide explicitly which regime you are in. Missing the target by an order of magnitude ends the question; landing within a factor of two means the estimate proved nothing and the exact numbers are now compulsory.
- Finish every allocation or capacity problem by listing how much of each constraint your answer consumes. The constraint you exhaust is the one that explains the answer, and the one you do not is where a per-unit heuristic will mislead you.
- Always verify a backwards solution by running it forwards. Attempts stay on this device, so the extra minute costs nothing and catches the inverted-fraction error that makes this family's most attractive wrong answer.

Glossary

Break-even point: The value at which two options cost or take the same. Below it one option wins and above it the other does, which is why a single quoted case never answers a comparison question.

Binding constraint: The limit that is fully consumed by the chosen solution. Identifying it explains the answer and reveals which per-unit ratios are relevant and which are noise.

Slack: The unused portion of a constraint at a proposed solution. Spare capacity in one resource is what lets a mix include items that look inefficient on the binding resource.

Bisection: Halving a search space with each test. It reduces a search of n candidates to about $\log$ base two of n tests, but only where the fault is monotone along the chain.

Working backwards: Solving from a known end state by inverting each operation in reverse order. The reliable method whenever the final value is given and every step is individually invertible.

Order-of-magnitude estimate: A rough calculation intended to settle feasibility rather than produce a figure. It is decisive when the answer is off by a factor of ten and useless when the margin is within a factor of two.

Root cause: The condition whose removal stops the fault recurring, as opposed to a symptom whose

treatment suppresses it temporarily. A candidate cause that cannot be tested by changing it remains a hypothesis.

Monotone fault: A fault where everything on one side of a break behaves differently from everything on the other. The precondition that makes bisection valid, and the assumption most often left unstated.

Study this next

- Problem solving skill hub: <https://learn.novusstreamsolutions.com/aptitude/skills/problem-solving>
- Problem solving practice bank: <https://learn.novusstreamsolutions.com/cognitive-skills/problem-solving>
- Diagrammatic reasoning lesson: <https://learn.novusstreamsolutions.com/aptitude/skills/diagrammatic-reasoning/lesson>
- Maintenance technician suite: <https://learn.novusstreamsolutions.com/aptitude/tests/maintenance-technician>
- Supply chain planner suite: <https://learn.novusstreamsolutions.com/aptitude/tests/supply-chain-planner>
- Problem solving lesson on Novus Learn: <https://learn.novusstreamsolutions.com/aptitude/skills/problem-solving/lesson>

Where this material comes from

- Every figure above was computed and checked for Novus Learn: the courier break-even at about 53.3 km, the six-test bisection bound for 40 candidates, the 4,400 starting budget verified forwards, and the four-X-two-Y loading enumerated across all feasible whole-pallet mixes.
- Terminology checked against the public Wikipedia articles 'Binary search algorithm', 'Break-even', 'Shadow price' and 'Root cause analysis'. Definitions only; every scenario above is original.
- Novus Learn aptitude construct registry (catalog seed) for the construct scope and related suite mapping.
- Public educational framing only: not affiliated with any official exam board, publisher or employer, and no copyrighted test item is reproduced.

Educational preparation only. Novus Learn does not administer official exams and does not guarantee scores or hiring outcomes.

Recommended preparation

- Build numerical reasoning: <https://learn.novusstreamsolutions.com/aptitude/skills/numerical-reasoning/lesson>
- Practice data interpretation: <https://learn.novusstreamsolutions.com/aptitude/skills/data-interpretation/lesson>
- Strengthen critical thinking: <https://learn.novusstreamsolutions.com/aptitude/skills/critical-thinking/lesson>

Answer keys, scoring and privacy

Answer keys and scoring logic stay server-side and are never included in any download or export.

Formal suite answer keys and scoring logic stay on the server and are not part of any download, in any format. Downloadable keys exist only for the open, untimed practice material (the practice packs on the puzzles, cognitive-skills and reasoning practice lab pages), where the answers are already public teaching content.

Novus Learn needs no account. Your practice history lives in this browser's storage on this device, and is sent to a server only if you create an account and switch on backup. This file contains no attempt, session or result link, so it is safe to share.

Private practice result. Not an official exam certificate, employer decision, hiring signal, admissions decision, or guaranteed outcome. Scores stay on this device unless you export them.

Answer keys and scoring logic stay server-side and are never included in any download or export.

Source: <https://learn.novusstreamsolutions.com/aptitude/tests/marketing-analyst/study-outline>