

Aircraft-maintenance aptitude: study guide

Aviation, aerospace, air traffic, and airports · suite apt-075-aircraft-maintenance-aptitude · generated 2026-09-15T15:39:52.885Z

Detail	Value
Title	Aircraft-maintenance aptitude: study guide
Generated	2026-09-15T15:39:52.885Z
Fixture/version	apt-075-aircraft-maintenance-aptitude
Sector	Aviation, aerospace, air traffic, and airports
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 Aircraft-maintenance aptitude

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-075-aircraft-maintenance-aptitude&mode=guided>
- Start mini-test:
<https://learn.novusstreamsolutions.com/aptitude/practice/new?bank=apt-075-aircraft-maintenance-aptitude&mode=mini>
- Start full simulation:
<https://learn.novusstreamsolutions.com/aptitude/practice/new?bank=apt-075-aircraft-maintenance-aptitude&mode=full>

Skills covered, in full

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

Mechanical comprehension

Forces, motion, gears, pulleys, levers, fluids, pressure, and basic machines.

Mechanical comprehension is the ability to predict what a physical system will do (which way a gear turns, how hard you have to pull, how far the load actually rises) from forces, moments and the simple machines, rather than from a memorised formula sheet. It carries real weight in apprenticeship entry batteries, military technical selection, and screening for maintenance, machine operation and rigging roles. The same reasoning is daily work on site: sizing a jack, choosing a block-and-tackle arrangement, deciding why a belt slips under load but not at idle. 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. Balance moments about a pivot (effort x effort arm = load x load arm) and name whether a given tool is a first-, second- or third-class lever.
2. Work out the speed, torque and rotation direction of a driven gear from tooth counts, including what an idler gear does and does not change.
3. Count the rope parts that support a moving block to get the mechanical advantage of a pulley system, then state how much rope must be pulled per metre of lift.
4. Apply pressure = force / area to a hydraulic jack and predict the piston-travel penalty that comes with the force gain.
5. Predict rpm across a belt or chain drive from pulley or sprocket diameters, and say when slip breaks the prediction.
6. State the trade-off that governs every simple machine: multiplied force always costs distance, and friction always makes the delivered advantage smaller than the ideal one.

Worked examples and pitfalls

Lever: balance the moments, then name the class: A wheelbarrow carries 60 kg whose centre of mass sits 0.4 m from the wheel axle, and you grip the handles 1.2 m from that same axle. Take moments about the axle: lift force x 1.2 m = 60 kgf x 0.4 m, so the lift force is $60 \times 0.4 / 1.2 = 20$ kgf, about 196 N. The mechanical advantage is simply the arm ratio, $1.2 / 0.4 = 3$. Two traps live in this item. The first is measuring the load arm from your hands instead of from the fulcrum; the arm is always the perpendicular distance to the pivot, and the pivot here is the wheel contact, not the barrow body. The second is calling it a first-class lever because that is the one everyone pictures. A wheelbarrow is second class: fulcrum at one end, load in the middle, effort at the far end, which is why its mechanical advantage is always greater than one. Compare tweezers or a pair of tongs, where the effort sits between fulcrum and load: that is third class, mechanical advantage below one, and you are deliberately trading force away to buy speed and control at the tip. Your own forearm lifting a weight is the same arrangement, which is why a 5 kg dumbbell loads the biceps far more than 5 kg.

Gear trains: tooth counts set speed, meshes set direction: A 12-tooth driver turning at 300 rpm meshes directly with a 36-tooth gear. The ratio is driven teeth over driver teeth, $36/12 = 3:1$, so the output turns at $300/3 = 100$ rpm and, ignoring friction, carries about three times the torque. One external mesh reverses rotation, so the output turns opposite to the driver. Now drop a 20-tooth idler between them. Step it through: $300 \times 12/20 = 180$ rpm at the idler, then $180 \times 20/36 = 100$ rpm at the output. The overall ratio is unchanged at 3:1, the idler's tooth count cancels, but there are now two external meshes, two reversals, so the output turns the SAME way as the driver. That is the whole reason idlers are fitted. The tempting wrong answer

treats the idler as another reduction stage and reports 180 rpm or some product of both ratios. The check that never fails: only the first and last gear in a simple train affect the ratio, and the direction depends on whether the number of external meshes is odd (reversed) or even (same). An internal or ring mesh, as in a planetary set, does not reverse at all.

Pulleys: count the rope parts that carry the load: A 200 kg load, about 1,962 N. Hung from a single pulley bolted to a beam, the pulley only changes the direction you pull; both rope parts still meet at a fixed axle, the mechanical advantage is 1, and you pull the full 200 kgf. Hang the pulley on the load instead, with one rope end anchored above and the other in your hands, and two rope parts now support the moving block: mechanical advantage 2, effort about 100 kgf or 981 N, but you must pull 2 m of rope for every 1 m the load rises. Build a tackle with four parts supporting the moving block and the effort falls to $200/4 = 50$ kgf, about 490 N, at the cost of 4 m of rope pulled per metre of lift. The trap is counting sheaves instead of supporting parts. A three-sheave arrangement gives three parts if the dead end is made off to the fixed block and four if it is made off to the moving block, and the answer differs by 33 percent. The second trap is treating the figure as delivered force: real sheaves lose a few percent each to bearing and rope friction, so quoted mechanical advantage is the ideal velocity ratio, and the effort you actually feel is higher.

Hydraulics: pressure is shared, force and travel are traded: A jack has a 2 square centimetre input piston and a 50 square centimetre output ram. Push the small piston with 100 N and the pressure in the fluid is $100 / 2 = 50$ N per square centimetre, which is 500 kPa. Pascal's principle says every part of the confined fluid sees that same 500 kPa, so the large ram feels $50 \text{ N/cm}^2 \times 50 \text{ cm}^2 = 2,500 \text{ N}$. Mechanical advantage 25. Nothing is free: fluid is effectively incompressible, so volume in equals volume out. A 25 cm stroke on the small piston moves $2 \times 25 = 50$ cubic centimetres, and 50 cubic centimetres spread over a 50 square centimetre ram raises it 1 cm. Twenty-five centimetres of pumping buys one centimetre of lift: exactly the force-times-distance bargain a lever strikes. Candidates lose this item by assuming both pistons travel the same distance, or by concluding the jack manufactures energy. A related item asks about a tank: the pressure at the bottom depends on the height of fluid above and its density, not on the width of the vessel, so a narrow 3 m standpipe reads the same bottom pressure as a 3 m deep swimming pool.

Belts and chains: same belt speed, different rpm: A motor pulley 100 mm in diameter runs at 1,450 rpm and drives a 250 mm pulley through an open V-belt. The quantity that is genuinely shared is belt speed, not rpm, so the driven pulley turns at $1,450 \times 100/250 = 580$ rpm. Note the direction of the ratio: with gears you divide by driven teeth, with belts you divide by driven diameter, and the arithmetic looks the same only because both are proportional to circumference. An open belt drives both pulleys the same way; a crossed belt reverses the driven pulley, which is the entire point of the classic crossed-belt diagram item. The trap is slip. A flat or V-belt under an overload slips, so the driven speed falls below 580 rpm while the motor holds 1,450: the symptom of a glazed belt or a weak tensioner. A chain or a toothed timing belt physically cannot slip a tooth without damage, which is why camshaft and indexing drives use them and why a belt-driven answer and a chain-driven answer to the same question can legitimately differ.

Springs in series and parallel: the answer most people reverse: Two identical coil springs, each 20 N/mm. Mount them side by side, both carrying the same load, and they act in parallel: stiffness adds to 40 N/mm, so a 400 N load compresses the pair $400/40 = 10$ mm. Now stack them end to end, in series. Each spring carries the whole 400 N, force is not split down a chain, so each deflects $400/20 = 20$ mm and the total is 40 mm. The combined stiffness is $400/40 = 10$ N/mm, half of one spring. Most candidates reverse the two because 'series' sounds like the case where things add. Two reliable checks: in parallel the arrangement is always stiffer than one spring, and in series it is always softer than the softest spring. That is also why suspension designers add a helper spring in parallel to raise rate under load, and why a long slender bolt clamps more forgivingly than a short stubby one carrying the same preload.

How to practise this skill

- Sketch the free body and mark the pivot before writing a single number. Mechanical items are lost by choosing the wrong fulcrum far more often than by mis-multiplying.
- Fix the three lever classes to a tool you can picture (pliers and a see-saw for first class, a wheelbarrow and a nutcracker for second, tweezers and your own forearm for third), and say the class out loud before answering.
- Sanity-check every answer against the work rule: if a machine multiplies force it must divide distance. An answer where the effort both moves less and is smaller than the load is wrong somewhere.
- Convert units before you multiply, not after. A moment computed with millimetres on one side and metres on the other is off by a thousand and looks perfectly plausible.
- For any direction question, walk the mesh or the belt with a finger: each external gear pair reverses, an internal or ring mesh does not, an open belt keeps direction and a crossed belt flips it.
- Work untimed until you can explain the reasoning aloud, then add the clock. Attempts are recorded on this device only, so a messy first pass through gear trains costs you nothing but time.

Glossary

Moment (torque): A force multiplied by the perpendicular distance from its line of action to the pivot, in newton-metres. A body is in rotational balance when clockwise moments equal anticlockwise moments about the same point.

Fulcrum: The pivot a lever turns about. Its position relative to the load and the effort defines the lever class and therefore whether the tool multiplies force or multiplies movement.

Mechanical advantage: Output force divided by input force. Ideal mechanical advantage is the geometric figure; actual mechanical advantage is what you get after friction, and it is always lower.

Velocity ratio: Distance moved by the effort divided by distance moved by the load. In a lossless machine it equals the ideal mechanical advantage, which is why force gain always costs travel.

Gear ratio: Driven teeth divided by driver teeth. A ratio above one reduces speed and multiplies torque; only the first and last gear of a simple train affect it.

Idler gear: A gear placed between driver and driven purely to add a mesh. It reverses the output direction and bridges a distance without altering the overall ratio.

Pascal's principle: Pressure applied to a confined fluid is transmitted undiminished to every part of it, which is what lets a small piston at high pressure produce a large force on a bigger ram.

Head: The height of a fluid column above a point. Pressure at depth is density x gravity x height and is independent of the vessel's width or shape.

Study this next

- Mechanical comprehension skill hub:
<https://learn.novusstreamsolutions.com/aptitude/skills/mechanical-comprehension>
- Millwright and industrial mechanic suite:
<https://learn.novusstreamsolutions.com/aptitude/tests/millwright-and-industrial-mechanic>
- Spatial reasoning lesson: <https://learn.novusstreamsolutions.com/aptitude/skills/spatial-reasoning/lesson>
- Skilled trades, maintenance and repair careers:
<https://learn.novusstreamsolutions.com/careers/families/skilled-trades-maintenance-and-repair>
- Encyclopedia search: simple machine:
<https://learn.novusstreamsolutions.com/search?q=Simple%20machine>
- Mechanical comprehension lesson on Novus Learn:
<https://learn.novusstreamsolutions.com/aptitude/skills/mechanical-comprehension/lesson>

Where this material comes from

- Worked examples written for Novus Learn from standard introductory mechanics: moments about a pivot, gear ratios from tooth counts, rope-part counting for tackle, and Pascal's principle.
- Terminology checked against the public Wikipedia articles 'Simple machine', 'Lever', 'Gear train' and 'Pascal's law'. Definitions only; every number and item 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.

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:

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.

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.

Data checking and clerical accuracy

Matching, filing, coding, record checking, and identifying discrepancies.

Clerical accuracy is the ability to handle a large volume of records correctly and consistently: filing them in the right order, coding them against a key, spotting duplicates that do not look like duplicates, and holding that standard on the two-hundredth record as well as the second. Where error detection asks whether two things match, clerical accuracy asks whether you can apply a rule reliably at volume. It is assessed in administrative, records, clinical admin, school office, evidence-handling and insurance operations selection, and it is exactly the skill an employer is buying when they hire for a data-heavy back-office role. 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. Apply an alphabetical filing rule consistently, and state whether a given system files letter-by-letter or word-by-word. The two produce different, equally correct orders.
2. Sort alphanumeric references correctly under both string ordering and natural numeric ordering, and explain why record systems zero-pad.
3. Code records against a written key, handling every boundary condition (under, up to, and over, inclusive ranges) exactly as written rather than as intuited.
4. Identify duplicate records that differ only in formatting (case, whitespace, date order, punctuation inside a reference), and route genuinely ambiguous ones to exceptions instead of resolving them by guess.
5. Measure your own accuracy decay across a long batch and use the measurement to set a realistic attempt rate.
6. Convert a stated scoring rule into a target items-per-minute figure, and show why the optimum rate moves when the penalty multiplier changes.

Worked examples and pitfalls

Word-by-word or letter-by-letter: two correct answers, one rule: File these four names: Van Dyke, Vandenberg, Van Horn, Vance. Under word-by-word filing, each space-separated unit is compared in turn and 'nothing files before something', so the first unit 'Van' sorts ahead of both 'Vance' and 'Vandenberg'. The order is Van Dyke, Van Horn, Vance, Vandenberg. Under letter-by-letter filing, spaces are ignored entirely, so the keys are VANCE, VANDENBERG, VANDYKE, VANHORN, and the order is Vance, Vandenberg, Van Dyke, Van Horn. Both orders are correct filing; only one is correct for the system you are working in. Test items state the convention in the instructions, and the candidates who lose marks are almost always the ones applying whichever convention their previous employer used. The same fork appears with prefixes and punctuation: whether Mc and Mac interfile, whether St is treated as Saint, whether a hyphen counts as a space. Read the rule, restate it to yourself in one sentence, then apply it mechanically and do not let a name that 'obviously' belongs somewhere override it.

Alphanumeric codes: A-102 at the front or the back of the drawer: Sort the references A-7, A-12, A-70, A-102, B-3. Under natural numeric ordering, the way a person reads them, the answer is A-7, A-12, A-70, A-102, B-3. Under plain string ordering, the way most software sorts by default, each character is compared in turn, so '1' comes before '7' and the answer is A-102, A-12, A-7, A-70, B-3, with A-102 first rather than last. Both are defensible; a filing item is testing which one the stated system uses. This is also why serious record systems zero-pad their references: rewrite the set as A-007, A-012, A-070, A-102 and the string sort and the numeric sort agree, permanently. If you are ever asked to design or clean a reference scheme, pad to a fixed width and the whole class of problem disappears. In a test, the giveaway is a set that deliberately mixes one-, two- and three-digit suffixes; that mixture exists only to separate candidates who apply the stated rule from candidates who apply the intuitive one.

Coding to a key: every mark is at a boundary: A claims routing key: under 500 pounds with no injury reported goes to CS1; under 500 with injury goes to CI2; 500 to 4,999 with no injury goes to CS3; 500 to 4,999 with injury goes to CI4; 5,000 and over goes to CE5 regardless of injury. Now code five records. R-118 at 499.99, no injury: CS1, since 499.99 is under 500. R-119 at exactly 500.00, no injury: CS3, because 'under 500' excludes 500 itself and the second band starts there. R-120 at 4,999.00 with injury: CI4, since the band is inclusive at its top. R-121 at exactly 5,000.00 with no injury: CE5, because the top band is 'and over' and its 'regardless of injury' clause overrides the injury split that governs the lower bands. R-122 at 86.40 with injury: CI2. Four of those five decisions turn on a boundary, and that is not an accident, coding items are written so that the interior cases are trivial and every discriminating mark sits on an edge. Before coding anything, underline the boundary words: under, up to, and over, between, inclusive. Then decide, once, what each one does to the endpoint, and apply that decision to every record in the batch.

Duplicates that are not textually identical: Three rows arrive in a merge. Row 1: SMITH, JANE | 07/04/1988 | AB123456C. Row 2: Smith, Jane | 1988-04-07 | AB 123456 C. Row 3: SMITH, JANE | 04/07/1988 |

AB123456C. Rows 1 and 2 are the same person: normalise case, strip the spaces from the reference, and convert the ISO date and they match exactly. Row 3 is the genuinely hard one. If the file is in day/month order it is a different date of birth and possibly a different person; if that row came from a system using month/day order it is the same record again. You cannot tell from the row itself, so the correct action is to send it to the exception queue with the ambiguity noted, not to merge it and not to discard it. This is the judgement that separates competent records work from the appearance of it: the goal is not to make every row disappear, it is to make every decision defensible. In test form the item usually asks 'how many distinct individuals are represented', and the answer is often given as a range or accompanied by a 'cannot determine' option for exactly this reason.

Where accuracy actually decays on a long batch: Run a self-measurement rather than trusting a number from anyone: take 200 record pairs, split them into four blocks of 50, and log errors per block. Most people find block 1 slightly worse than block 2, a warm-up cost, and then a rise across blocks 3 and 4 as vigilance falls. The reason to measure it is that the arithmetic of small percentages is brutal at volume. Checking 200 pairs at 98 percent accuracy passes 4 bad records; at 99.5 percent it passes 1. Scale that to a realistic month of 20,000 records and the same two accuracy rates mean 400 defects against 100: a four-fold difference in downstream rework from a 1.5 point difference that would look like noise on a single test. Once you know where your own curve turns, the intervention is cheap: a deliberate twenty-second break at that point, or splitting the batch so the hardest records fall in your strongest block. Practice sessions on this site are recorded on your own device, so building a block-by-block picture across several sessions costs nothing but the logging.

Setting an attempt rate from the scoring rule: A 120-item checking test with a 10-minute limit, scored as correct minus incorrect. Suppose practice has told you that you hold 92 percent at ten items a minute and 96 percent at seven and a half. Fast: $10 \times 10 = 100$ attempted, 92 correct and 8 wrong, score 84. Careful: $10 \times 7.5 = 75$ attempted, 72 correct and 3 wrong, score 69. Fast wins by 15. Now change the rule to correct minus three times incorrect: fast scores $92 - 24 = 68$, careful scores $72 - 9 = 63$, and the 15-point gap has shrunk to 5. Now suppose your real accuracy at ten a minute is 80 percent rather than 92. A gap most people do not discover until they measure it. Fast now gets 80 right and 20 wrong: under simple correct-minus-incorrect that is 60, already behind the careful strategy's 69, and under the triple penalty it is $80 - 60 = 20$ against 63. Same test, same person, opposite advice, and the two deciding variables are the penalty multiplier, which the instructions hand you, and your own accuracy-at-speed, which only measurement gives you. Do not pick a pace from temperament. Measure two rates in practice, write both accuracy figures down, and do this arithmetic before the test rather than during it.

How to practise this skill

- Before the first record of any batch, write the filing or coding rule at the top of the page in your own words. The single largest source of clerical error is applying a remembered rule from a previous system.
- Underline the boundary words in a coding key (under, up to, and over, inclusive), and resolve each endpoint once. Interior cases carry almost no marks; the edges carry nearly all of them.
- Normalise before you compare: mentally strip case, spaces and punctuation from references, and restate dates in one fixed order. Half of apparent duplicates are formatting, and half of apparent non-duplicates are too.
- When a record is genuinely ambiguous, mark it and move on. Time spent resolving one unresolvable row is taken from thirty rows you could have done correctly, and a guessed merge is worse than a flagged one.
- Measure your accuracy at two different speeds in practice and write both numbers down. You cannot choose an attempt rate rationally without them, and the fast rate is almost never as accurate as it feels.
- Practise on batches long enough to hit your own fatigue point, not on ten-item samples. The construct is specifically about sustained accuracy, and a ten-item drill measures the part of the curve that was never in doubt.

Glossary

Word-by-word filing: An alphabetical convention that compares space-separated units in turn, with a shorter first unit filing before a longer one. Under it, Van Horn files before Vance.

Letter-by-letter filing: An alphabetical convention that ignores spaces and punctuation entirely, comparing the run of letters. Under it, Vance files before Van Dyke.

Natural sort: Ordering that reads embedded digit runs as numbers, so A-7 precedes A-12. Contrasted with string ordering, which compares characters one at a time and puts A-102 first.

Zero padding: Writing references to a fixed width by adding leading zeros (A-007 rather than A-7) so that string ordering and numeric ordering produce the same result. The standard structural fix for alphanumeric filing errors.

Primary and secondary sort key: The field sorted on first and the field used to break ties within it. A batch sorted by site then by surname will look wrong if the two keys are applied in the opposite order.

Canonicalisation: Converting records to a single standard form (one case, no stray whitespace, one date order, punctuation stripped from references) before any comparison, so that formatting differences stop masquerading as data differences.

Exception queue: The destination for records that cannot be resolved from the information available. Routing an ambiguous record here is a correct outcome; guessing it into a merge is not.

Boundary condition: The endpoint of a coded band, where wording such as under, up to or and over decides which side a value falls. Coding tests concentrate their discriminating items here.

Study this next

- Data checking and clerical accuracy skill hub:
<https://learn.novusstreamsolutions.com/aptitude/skills/clerical-accuracy>
- Clerical accuracy practice in the Cognitive Skills Lab:
<https://learn.novusstreamsolutions.com/cognitive-skills/clerical-accuracy>
- Error detection lesson: <https://learn.novusstreamsolutions.com/aptitude/skills/error-detection/lesson>
- Administrative and clerical entry suite:
<https://learn.novusstreamsolutions.com/aptitude/tests/administrative-and-clerical-entry>
- Clerical checking suite: <https://learn.novusstreamsolutions.com/aptitude/tests/clerical-checking>
- Data checking and clerical accuracy lesson on Novus Learn:
<https://learn.novusstreamsolutions.com/aptitude/skills/clerical-accuracy/lesson>

Where this material comes from

- All filing sets, reference codes, routing keys and records above were invented for this lesson. The two filing orders, the two sort orders and every score calculation in the attempt-rate example were worked through and checked by recomputation.
- Filing and sorting conventions cross-checked against standard public descriptions such as the Wikipedia articles 'Alphabetical order', 'Natural sort order' and 'Collation'. Terminology only; the examples are original.
- 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.

Situational judgment

Evaluating workplace responses against role-relevant principles.

Learn a repeatable way to compare workplace responses: establish the facts, identify duties and risks, respect role boundaries, then choose a proportionate first action. This is educational preparation, not an official scoring guide.

What you should be able to do after this lesson:

1. Separate facts stated in a scenario from assumptions that the scenario does not support.
2. Rank response options by immediate risk, policy or role obligations, proportionality, and follow-through.
3. Explain why a strong first action is better than passive, punitive, or unauthorized alternatives.

Worked examples and pitfalls

Worked scenario: an unverified safety concern: A colleague reports a possible equipment fault while a deadline is approaching. First distinguish the known fact, the report, from the unverified cause. A strong response protects people and affected work, checks the concern through the right channel, tells the relevant lead, and records what was done. Ignoring the report underreacts; shutting down unrelated work or accusing someone before checking the facts overreacts.

Method: facts, duties, risks, response: Write four short notes before ranking options: what is known, who may be affected, which duty or boundary applies, and what safe next step is available now. Prefer an action that addresses the immediate issue and creates useful follow-through. Do not reward an option merely because it sounds decisive.

How to practise this skill

- Answer the question asked: best first action, worst action, or complete response are different tasks.
- Check whether an option acts within the person's authority and escalates only as far as the risk requires.
- When two options look reasonable, prefer the one that gathers missing facts and communicates ownership.

Glossary

Proportionality: Matching the urgency and scope of a response to the evidence, likely impact, and authority available.

Role boundary: The limit of what a person may decide or do without approval, specialist help, or escalation.

Follow-through: Confirming ownership, recording the decision, and checking that the issue was actually resolved.

Study this next

- Situational judgement test:
<https://learn.novusstreamsolutions.com/search?q=Situational%20judgement%20test>
- Workplace ethics: <https://learn.novusstreamsolutions.com/search?q=Workplace%20ethics>
- Decision-making: <https://learn.novusstreamsolutions.com/search?q=Decision-making>
- Situational judgment lesson on Novus Learn:
<https://learn.novusstreamsolutions.com/aptitude/skills/situational-judgement/lesson>

Where this material comes from

- Novus Learn original situational-judgment suite scenarios and published construct mapping.
- Novus educational framework: facts, duties, risks, role boundaries, proportional action, and follow-through.

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

Recommended preparation

- Build mechanical comprehension:
<https://learn.novusstreamsolutions.com/aptitude/skills/mechanical-comprehension/lesson>
- Build numerical reasoning:
<https://learn.novusstreamsolutions.com/aptitude/skills/numerical-reasoning/lesson>
- Build 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/aircraft-maintenance-aptitude/study-outline>