

Research and fact checking: study guide

Media, communications, language, and creative operations · suite apt-340-research-and-fact-checking · generated 2026-09-15T17:47:33.176Z

Detail	Value
Title	Research and fact checking: study guide
Generated	2026-09-15T17:47:33.176Z
Fixture/version	apt-340-research-and-fact-checking
Sector	Media, communications, language, and creative operations
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 Research and fact checking

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-340-research-and-fact-checking&mode=guided>
- Start mini-test:
<https://learn.novusstreamsolutions.com/aptitude/practice/new?bank=apt-340-research-and-fact-checking&mode=mini>
- Start full simulation:
<https://learn.novusstreamsolutions.com/aptitude/practice/new?bank=apt-340-research-and-fact-checking&mode=full>

Skills covered, in full

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

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.

Reading comprehension

Understanding passages, instructions, policies, technical material, and workplace documents.

Reading comprehension is the ability to take a passage, a policy clause, a procedure or a technical note and answer accurately about what it says, what it means, and how it applies to a case in front of you. It is assessed directly in casework, policy analysis, immigration and language-focused batteries, and it sits underneath almost every other written test as the thing that has to work before anything else can. On the job it is the difference between applying an eligibility rule correctly the first time and generating an appeal. Everything you practise here is stored on this device; there is no account and no upload.

What you should be able to do after this lesson:

1. Choose a main-idea answer by testing scope first, and reject options that are true but too narrow, too broad, or about a different passage entirely.
2. Apply a written rule to specific cases, reading AND, OR, 'at least', 'more than' and 'unless' exactly as written and checking the exception clause last.
3. Infer the meaning of a word from the sentence around it rather than from its most common everyday sense.
4. Track referents across sentences - it, this, the former, the latter, the team - and name what each one points at.
5. Separate the author's own stance from views the author is reporting, and read hedging as evidence of stance.
6. Work through a numbered procedure with conditional steps, including cases where two steps both apply.

Worked examples and pitfalls

Main idea: scope, not truth: Passage: 'When the Kirkwall depot adopted route-optimisation software in 2021, its drivers covered 9 percent fewer kilometres in the first quarter, and fuel spending fell by a similar margin. Delivery times improved on urban routes only; rural routes were unchanged, and two drivers reported longer split shifts. The depot manager has kept the software but now overrides its rural schedules by hand.' Which option states the main idea? (a) 'Route-optimisation software reduces fuel costs.' True, and it is exactly one third of the passage - it drops the mixed results and the override, so it is too narrow. (b) 'The software failed at Kirkwall.' Contradicted: mileage and fuel both improved, and the depot kept it. (c) 'Rural deliveries are harder to optimise than urban ones.' A generalisation the passage never makes; it reports one depot, one year, without explaining why rural routes were unchanged. Too broad. (d) 'At one depot the software produced clear mileage and fuel savings but uneven results elsewhere, so it is now used with manual override.' Correct - it covers all three sentences and no more. The lesson generalises: main-idea distractors are usually wrong for scope, and the two commonest failures are a true detail promoted to the main point, and a reasonable-sounding claim the passage never actually makes.

Reading a clause exactly: and, at least, more than, unless: Rule: 'An employee may claim the remote-working allowance if they work from home on at least three scheduled days per week AND their home is more than 40 km from their assigned office, unless they already receive the travel-card subsidy.' Four cases. Priya works from home four days, lives 52 km away, and receives the travel-card subsidy: not eligible - she satisfies both qualifying conditions, and the 'unless' clause overrides both. Tom works from home three days, lives 38 km away, no subsidy: not eligible, because the connective is AND and 38 is not more than 40. Dan works from home two days, lives 60 km away, no subsidy: not eligible, failing the days condition for the same structural reason. Mira works three days, lives 41 km away, no subsidy: eligible - 'at least three' includes exactly three, and 41 is more than 40. Three separate traps live in one sentence: reading AND as OR under time pressure, treating 'more than 40' as including 40, and forgetting that an exception clause outranks the conditions it follows. The reliable method is to rewrite the rule as a checklist - condition 1, connective, condition 2, then exception - before you look at any case.

Vocabulary in context: the evidence is in the next clause: Sentence: 'The auditor's remarks were characteristically dry, and the board took a full minute to realise she had been criticising them.' Which meaning does 'dry' carry here? (a) arid or lacking moisture, (b) dull and lacking interest, (c) understated and quietly ironic, (d) free of alcohol. Option (b) is the distractor that catches most people, because it is the commonest figurative sense and it half-fits an auditor. But the clause after 'and' is doing the work: a remark that takes a minute to land as criticism is understated, not boring - dullness would not delay recognition, it would only reduce attention. The answer is (c). The general method for vocabulary-in-context items is to ignore the word for a moment, read the rest of the sentence as evidence, and ask what property the sentence requires the word to have. Test items are built so that the most frequent sense of the word is available as an option and is wrong; if the most obvious meaning were correct, the item would measure nothing.

Referents: the former, the latter, and the noun three sentences back: Passage: 'Two remedies were proposed: a fixed surcharge on late filings, and a sliding penalty tied to the amount owed. The committee rejected the former on fairness grounds, noting that it would fall hardest on the smallest filers.' Which remedy was rejected? 'The former' is the first item mentioned, so the fixed surcharge. Candidates who skim choose the sliding penalty because it sits nearer to the pronoun, which is exactly why the item is written this way. Notice also the free consistency check built into the sentence: a fixed amount is the one that lands hardest on small filers, because the same sum is a larger share of a smaller liability, while a penalty tied to the amount owed scales with size by construction. When a passage gives you the reason alongside the reference, use it to confirm the reference. The harder version of this item type replaces 'the former' with a bare 'it' after a sentence containing three noun phrases, and the technique is the same: write the candidate referents in the margin and substitute each one into the sentence to see which produces a sentence the

passage could have meant.

Whose view is this? Stance versus report: Passage: 'Supporters of the levy argue that it will fund three new depots within a decade. That projection rests on traffic volumes holding at 2019 levels, which no forecast in the sector now expects.' Question: what is the author's position? (a) The author supports the levy. (b) The author opposes the levy. (c) The author doubts the depot projection. (d) The author has no view. The correct answer is (c). The stance markers are 'rests on', which frames the projection as dependent on a condition, and 'which no forecast now expects', which tells you the condition is not met. Option (b) is the trap, and it is a scope error dressed as a tone question: undermining one argument made by supporters is not opposition to the policy, and the passage says nothing about the levy's other merits or costs. Reading for stance means reading the author's verbs and qualifiers rather than the content of the views being described - a passage can spend four sentences laying out a position it goes on to dismantle in the fifth.

Procedures: when two steps both fire: Instructions: '1. Record the meter reading. 2. If the reading is lower than last month's, do not submit; raise a query instead. 3. Otherwise, submit the reading and file the photo. 4. If the reading is more than double last month's, submit the reading and flag it for review.' Case A: last month 4,180, this month 9,020. Step 2 does not apply, since 9,020 is higher. Step 3 applies: submit and file the photo. Does step 4 also apply? Double 4,180 is 8,360, and 9,020 is more than that, so yes - submit, file the photo, and flag for review. Steps 3 and 4 are not alternatives; nothing in the list makes them exclusive, and both instruct you to submit, which is consistent. Case B: last month 4,180, this month 4,090. Step 2 fires: do not submit, raise a query. Steps 3 and 4 never come into play, because step 3 begins 'otherwise' and step 4's condition is not met either. The trap in case A is treating a numbered list as a decision tree where exactly one branch executes. Read each step's condition independently unless the procedure explicitly says to stop.

How to practise this skill

- Map the passage before you answer: read the first and last sentence of each paragraph and write a three-word label for each. On a 400-word passage this takes about twenty seconds and makes every detail question a lookup instead of a search.
- For every detail answer, put a finger on the sentence that supports it. If you cannot point at one, you are answering an inference question by feel, and that is where the marks go.
- On main-idea items, test scope before truth. Ask of each option: does it cover the whole passage, and does it cover nothing outside it? Two options will usually be true and only one will be the right size.
- Rewrite any policy or eligibility clause as a numbered checklist with the connective and the exception written out separately. Assessors build these items around AND read as OR, and around the boundary values on 'at least' and 'more than'.
- Do not pre-read the options on inference and stance items. Answer in your own words first, then find the option that matches; reading four polished options first makes three of them sound reasonable.
- Log wrong answers by cause - scope, connective, referent, stance, boundary value - rather than by passage topic. The topic never repeats; the cause always does.

Glossary

Main idea: The claim that covers the whole passage and nothing beyond it. It is not the first sentence, not the most interesting fact, and not the most general statement available.

Scope: How much an answer option claims relative to the passage. Most wrong main-idea and inference options are true statements that are simply too narrow or too broad, which is why checking truth alone does not separate them.

Referent: The noun a pronoun or phrase such as it, this, the former, or the team points back to. Ambiguous referents are deliberate in test passages and are resolved by substituting each candidate into the sentence.

Connective: A logical joining word in a rule - and, or, unless, provided that, except where. AND requires

every condition; OR requires one; unless introduces an override that outranks what precedes it.

Boundary value: The exact number at the edge of a condition. 'At least three' includes three; 'more than 40' excludes 40; 'up to five' usually includes five. Items are written on these edges because that is where careless reading shows.

Stance: The author's own position, signalled by qualifiers, framing verbs and concessions rather than by direct statement. Distinct from any view the author reports.

Skimming versus scanning: Skimming is a fast first pass for structure and gist; scanning is a targeted hunt for a specific word or figure. Comprehension sections reward doing the first once and the second repeatedly.

Topic sentence: The sentence carrying a paragraph's central claim, most often first or last. Locating them across paragraphs gives you the passage's argument in about a fifth of the reading time.

Study this next

- Reading comprehension skill hub:
<https://learn.novusstreamsolutions.com/aptitude/skills/reading-comprehension>
- Practise the reading comprehension bank:
<https://learn.novusstreamsolutions.com/cognitive-skills/reading-comprehension>
- Verbal reasoning lesson: <https://learn.novusstreamsolutions.com/aptitude/skills/verbal-reasoning/lesson>
- Policy and program analysis suite:
<https://learn.novusstreamsolutions.com/aptitude/tests/policy-and-program-analysis>
- Immigration and asylum casework suite:
<https://learn.novusstreamsolutions.com/aptitude/tests/immigration-and-asylum-casework>
- Reading comprehension lesson on Novus Learn:
<https://learn.novusstreamsolutions.com/aptitude/skills/reading-comprehension/lesson>

Where this material comes from

- Worked items written for Novus Learn. The passages, eligibility rule, instruction list and figures above are original and invented for this lesson; no published or copyrighted test item is reproduced.
- Terminology follows standard, widely published usage in reading instruction and assessment writing - main idea, scope, referent, topic sentence.
- 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.

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:

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.

Deductive reasoning

Applying stated rules and constraints to reach valid conclusions.

Deductive reasoning is the disciplined application of rules you have been handed to a case you have been handed, stopping precisely where the rules run out. Where logical reasoning tests the machinery of inference, deductive items test restraint: they hand you a policy, a clue set or a specification and reward the candidate who reaches the one conclusion that is forced and declines the three that are merely plausible. It is the dominant construct in legal and policy admissions tests, critical-thinking style batteries with a cannot-say option, drone and aviation rule-knowledge screening, and any role where a written procedure has to be applied to an awkward real case. Practice stays on this device unless you export it.

What you should be able to do after this lesson:

1. Apply a multi-clause eligibility rule to a specific case and identify exactly which clause decides it, including when the case fails on one clause while satisfying every other.
2. Distinguish validity from truth, and explain why an assessment item requires you to accept a premise you believe to be false.
3. Choose the cannot-say or insufficient-information response correctly, and articulate the specific fact that would have been needed to decide.
4. Solve a small elimination grid from a clue set, showing which clue eliminated which cell.
5. Handle exclusive and inclusive or correctly, and know when eliminating one disjunct establishes the other.
6. Derive a bounded answer from numeric constraints (a maximum, a minimum, or a short list of possibilities) rather than guessing a single value the constraints do not fix.

Worked examples and pitfalls

A three-clause rule and the case that fails on one: A dispatch policy: an order ships the same day if it is placed before 14:00 local time, AND every line on the order is in stock, AND the account is not on credit hold. Three cases. Case A is placed at 13:40, has one of four lines backordered, and the account is clear. Same-day? No. The word AND makes this a conjunction: all three clauses must hold, and one failure is fatal regardless of how comfortably the others pass. The reason candidates get this wrong is not logic, it is arithmetic instinct. Two out of three feels like a pass, and the 13:40 timestamp is written to feel like the clause the item is really about. Case B is placed at 14:00 exactly, fully in stock, no hold. Same-day? No, because before 14:00 excludes 14:00 itself. Boundary wording is deliberate in policy items; before, by, up to and no later than are four different lines. Case C is placed at 09:15, fully in stock, account on credit hold, and the

customer has paid the invoice that caused the hold but the hold has not been lifted in the system. Same-day? No. The rule keys on the account being on hold, not on whether it deserves to be, and importing your judgement about what the policy is trying to achieve is the single most reliable way to fail this construct.

Valid but false, true but unjustified: All mammals lay eggs. A whale is a mammal. Therefore a whale lays eggs. The argument is VALID, the conclusion follows inescapably from the premises, and its first premise is false, so the argument is unsound and the conclusion is false. Validity is a property of the shape, soundness is validity plus true premises. Now the mirror image. Some birds cannot fly. A penguin is a bird. Therefore a penguin cannot fly. The conclusion is true, but the argument is invalid: some birds cannot fly leaves open which ones, so the premises do not force it. A true conclusion reached by an invalid route earns nothing. This matters practically because deductive assessment items routinely use premises that are false in the real world, all deliveries to the northern depot arrive by rail, precisely to test whether you reason from the page or from memory. The instruction to treat the statements as true is not a formality; a candidate who silently corrects a premise is answering a different question. Read the premise, accept it for the duration of the item, and put your world knowledge down.

When cannot say is the correct answer: Given: every technician on the night shift holds a confined-space ticket. Priya holds a confined-space ticket. Candidate conclusion: Priya works the night shift. The answer is cannot say, and the reason is worth naming precisely. The rule runs from night-shift technician to ticket-holder, and Priya has been placed at the ticket-holder end. Nothing in the premises says ticket-holders are exclusively night-shift technicians, so Priya could be a day-shift technician, a contractor or a safety officer. The fact that would have decided it is a converse or an exclusivity clause, only night-shift technicians hold the ticket, and being able to state that missing fact is the difference between guessing cannot-say and knowing it. Two calibration checks. First, cannot say is not a hedge for items you find hard; if you can construct one world where the conclusion is true and one where it is false while both satisfy every premise, cannot say is provably right. Second, cannot say is wrong whenever a contrapositive settles the matter: given the same premise and the fact that Sam does NOT hold a confined-space ticket, Sam definitely does not work the night shift, and answering cannot say there is a real error.

An elimination grid, solved clue by clue: Three analysts (Rowan, Sana, Theo) are each based in one of Cardiff, Dublin or Edinburgh and each covers one of energy, retail or transport. Clues: (1) the Dublin analyst covers retail; (2) Theo is not in Dublin and does not cover transport; (3) Sana is not based in Cardiff; (4) Rowan is not based in Edinburgh; (5) Sana does not cover transport. Work the cities first. Clues (2), (3) and (4) leave only two city arrangements: Rowan-Cardiff, Sana-Dublin, Theo-Edinburgh; or Rowan-Dublin, Sana-Edinburgh, Theo-Cardiff. Take the second. Rowan in Dublin covers retail by (1); Theo cannot cover transport by (2), so Theo covers energy; that leaves transport for Sana, which clue (5) forbids. The whole arrangement dies. So the first stands: Rowan in Cardiff, Sana in Dublin, Theo in Edinburgh. Now the subjects. Sana in Dublin covers retail by (1). Theo is not transport by (2), so Theo covers energy. Rowan takes transport, and clue (5) is satisfied because Sana has retail. Final grid: Rowan, Cardiff, transport; Sana, Dublin, retail; Theo, Edinburgh, energy. Two working habits do most of the lifting. Record negatives as well as positives. Theo-not-Dublin is as useful as any positive placement. And when a clue only bites after several others have landed, as clue (5) does here, park it and come back rather than abandoning the branch.

Or, and the two meanings it carries: A fault report states: either the sensor or the relay has failed, but not both. The relay tests good. Conclusion: the sensor failed. Valid, and comfortably so. The not both makes this an exclusive or, and eliminating one disjunct forces the other. Now strip the qualifier: either the sensor or the relay has failed. The relay tests good. Conclusion: the sensor failed. Still valid. Disjunctive syllogism works for the inclusive reading too, because inclusive or asserts at least one. The difference shows up on the OTHER side of the inference. With the exclusive version, learning the relay HAS failed tells you the sensor did not. With the inclusive version, learning the relay has failed tells you nothing at all about the sensor. Both may have gone. Candidates who treat every or as exclusive will confidently clear a good component; candidates who treat every or as inclusive will miss a genuine deduction. Read for the qualifier: but not both,

exactly one of, and either-but-not signal exclusivity, while a bare or in a technical document is normally inclusive. When a policy says applicants must hold a degree or five years of experience, the inclusive reading is almost certainly intended, and an applicant with both is not disqualified.

Numeric constraints give you a bound, not a value: Four people share a 28-night on-call rota. Each person takes at least four nights. Ama takes exactly twice as many nights as Bo. Cal takes eight. What is the largest number of nights Ama can take? Set it up: $A + B + C + D = 28$ with $C = 8$ and $A = 2B$, so $2B + B + 8 + D = 28$, giving $3B + D = 20$. Both B and D must be at least four, so $3B$ is at most 16 and B is at most 5 (since B must be a whole number and 3 times 6 is already 18, leaving $D = 2$). B is also at least four, so B is 4 or 5. B = 4 gives $A = 8$ and $D = 8$: the rota is 8, 4, 8, 8, totalling 28, all at least four: legal. B = 5 gives $A = 10$ and $D = 4$: the rota is 10, 5, 8, 5, totalling 28: also legal. So the answer is 10. The trap answer is 12, produced by checking the total and forgetting that D would then be 2, below the floor. Two lessons live here. The constraints do not fix a single rota, and a candidate who insists on one answer for every person will invent a value; the honest output is a bound plus the small set of legal arrangements. And the binding constraint is the floor on D, not the total, identifying which constraint actually stops you going further is the entire skill in this item type.

How to practise this skill

- Before touching a rule-application item, list the clauses on separate lines and mark each one as satisfied, failed or unknown for the case in front of you. Conjunctions are decided by the first failure, and writing them out stops you from being led by the clause the item made most vivid.
- Circle every boundary word, before, by, at least, no later than, more than, up to and including. A large share of deductive marks turn on whether the boundary value itself is inside or outside the rule.
- When you reach for cannot say, write the one sentence that would have settled the item. If you cannot name that missing fact, you are probably guessing, and if you can, you are almost certainly right.
- Practise accepting false premises out loud. Say the premise is false in the real world and I am reasoning from it anyway before you answer; the habit protects you from the items that are deliberately built on untrue statements.
- On grid puzzles, record negatives in the grid, not just confirmations, and note which clue produced each elimination so you can retrace a branch instead of restarting it.
- Finish every numeric-constraint item by substituting your answer back into all constraints, including the ones you did not use. Attempts stay on this device, so the extra thirty seconds of checking costs nothing and catches the off-by-one that the answer options are designed to reward.

Glossary

Valid argument: An argument whose conclusion cannot be false while its premises are true. Validity is about structure alone and says nothing about whether the premises are actually correct.

Sound argument: A valid argument whose premises are also true. Assessment items test validity; soundness is usually outside the scope you are given.

Premise: A statement offered as given. In deductive items you are instructed to treat premises as true for the duration of the question, even where they contradict your own knowledge.

Cannot say: The response indicating that the premises are consistent with the conclusion being either true or false. It is a positive finding about the evidence, not an admission of uncertainty.

Disjunctive syllogism: The valid inference from either P or Q together with not-P to the conclusion Q. It holds for both the inclusive and the exclusive reading of or.

Exclusive or: A disjunction asserting that exactly one alternative holds, signalled by but not both or exactly one of. Unlike inclusive or, it also lets you infer not-Q from P.

Counterexample: A single case satisfying every premise while making the conclusion false. Producing one is

the fastest proof that an argument is invalid or that an answer option is not forced.

Binding constraint: The constraint that actually stops a quantity going further. Naming it converts a vague numeric item into a definite maximum or minimum.

Study this next

- Deductive reasoning skill hub: <https://learn.novusstreamsolutions.com/aptitude/skills/deductive-reasoning>
- Deductive reasoning practice bank:
<https://learn.novusstreamsolutions.com/cognitive-skills/deductive-reasoning>
- Logical reasoning lesson: <https://learn.novusstreamsolutions.com/aptitude/skills/logical-reasoning/lesson>
- Legal reasoning suite: <https://learn.novusstreamsolutions.com/aptitude/tests/legal-reasoning>
- Critical-thinking style familiarization suite:
<https://learn.novusstreamsolutions.com/aptitude/tests/watson-glaser-style-critical-thinking-familiarization>
- Deductive reasoning lesson on Novus Learn:
<https://learn.novusstreamsolutions.com/aptitude/skills/deductive-reasoning/lesson>

Where this material comes from

- Worked items written for Novus Learn from standard introductory logic and constraint reasoning: conjunctive rule application, validity versus soundness, disjunctive syllogism, and elimination grids. The grid and the rota problem were each solved exhaustively and checked for a unique answer.
- Terminology checked against the public Wikipedia articles 'Validity (logic)', 'Soundness', 'Disjunctive syllogism' and 'Exclusive or'. Definitions only; every case and clue set 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

- Strengthen critical thinking: <https://learn.novusstreamsolutions.com/aptitude/skills/critical-thinking/lesson>
- Practice reading comprehension:
<https://learn.novusstreamsolutions.com/aptitude/skills/reading-comprehension/lesson>
- Practice verbal reasoning: <https://learn.novusstreamsolutions.com/aptitude/skills/verbal-reasoning/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/research-and-fact-checking/study-outline>