

Contract review: study guide

Legal, compliance, courts, and governance · suite apt-255-contract-review · generated
2026-09-15T15:27:47.987Z

Detail	Value
Title	Contract review: study guide
Generated	2026-09-15T15:27:47.987Z
Fixture/version	apt-255-contract-review
Sector	Legal, compliance, courts, and governance
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 Contract review

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-255-contract-review&mode=guided>
- Start mini-test:
<https://learn.novusstreamsolutions.com/aptitude/practice/new?bank=apt-255-contract-review&mode=mini>
- Start full simulation:
<https://learn.novusstreamsolutions.com/aptitude/practice/new?bank=apt-255-contract-review&mode=full>

Skills covered, in full

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

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:

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.

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.

Error detection

Comparing strings, records, transactions, forms, and data for mistakes.

Error detection is the skill of holding a source and a copy side by side and finding the one character, digit or field that moved, and of knowing which errors a given check will and will not catch. It is assessed directly in clerical checking, banking operations, records, proofreading and dispatch selection, and it underpins quality control anywhere a record is rekeyed or transcribed. The useful part of the skill is not staring harder; it is a repeatable scan procedure plus a set of arithmetic checks that catch what the eye misses. Everything you practise is stored on this device only, with no account and no upload unless you export it.

What you should be able to do after this lesson:

1. Name the four transcription error families (substitution, transposition, omission, duplication), and state

which of them a plain digit-sum check will fail to detect.

2. Run a fixed field-order comparison between a source record and a copy, reporting the specific field and character position that differs rather than a general impression.
3. Compute a modulus 11 check digit (ISBN-10 style) by hand and use it to decide whether a single record is internally consistent.
4. Apply the Luhn algorithm to a candidate account number, and state the one transposition case Luhn provably cannot catch.
5. Work out, from a stated scoring rule, whether guessing on an uncertain item has positive, zero or negative expected value.
6. Reconcile a document at batch level (line totals, subtotal, tax, grand total) and explain why an internally consistent document can still be wrong.

Worked examples and pitfalls

The four transcription error families: Source reference: INV-2024-08871. Four corrupted copies, one of each family. INV-2024-08571 is a substitution: a single character changed, 8 to 5. INV-2024-0871 is an omission: one character dropped, and the string is now a character short. INV-2024-088711 is a duplication: a character repeated, one character long. INV-2024-08817 is a transposition: the final 7 and 1 have swapped places, and crucially the string is the same length and contains exactly the same characters. That last property is why transposition is the hardest of the four to see and the most dangerous in practice. Anything that checks length catches omission and duplication. Anything that compares character sets or sums the digits catches substitution but not transposition, because $0+8+8+7+1$ and $0+8+8+1+7$ both come to 24. Reading the reference aloud catches substitution reliably and transposition poorly, because the ear is comparing sounds and both versions sound similar at speed. The only defences against transposition are a positional comparison and a weighted check digit.

A fixed scan order finds the one field that moved: Source record: Name Priyanka Ramaswamy / Account 4471-9026-3358 / Date of birth 14/03/1991 / Postcode SW1A 2AA / Phone 0161 496 0872. Copy: Name Priyanka Ramaswamy / Account 4471-9026-3358 / Date of birth 14/03/1991 / Postcode SW1A 2AA / Phone 0161 469 0872. The difference is in the phone field, where 496 has become 469: a transposition inside the middle block. Most candidates find it eventually; the ones who find it in eight seconds are the ones running a procedure. Always compare in the same field order, top to bottom, never skipping a field because it 'looks fine'. Compare long strings in the printed blocks rather than as a single run (4471, then 9026, then 3358) because short-term memory holds about four items and a twelve-digit run does not fit. Say the block silently, look, compare, move on. On a same-or-different item you may stop at the first mismatch; on a 'how many fields differ' item you must complete every field, and the instruction wording is what tells you which regime you are in. Getting this wrong in either direction costs marks: stopping early on a count item, or exhaustively checking a same-or-different item you had already resolved.

Modulus 11: why a weighted check digit catches a transposition: The ISBN-10 scheme multiplies the ten digits by descending weights 10, 9, 8 down to 1 and requires the total to be divisible by 11. Take 0-306-40615-2. Working left to right: $10 \times 0 = 0$, $9 \times 3 = 27$, $8 \times 0 = 0$, $7 \times 6 = 42$, $6 \times 4 = 24$, $5 \times 0 = 0$, $4 \times 6 = 24$, $3 \times 1 = 3$, $2 \times 5 = 10$, and the check digit $1 \times 2 = 2$. The sum is $0 + 27 + 0 + 42 + 24 + 0 + 24 + 3 + 10 + 2 = 132$, and $132 = 12 \times 11$ exactly, so the number is valid. Now transpose the 1 and the 5 to give 0-306-40651-2. The digits are identical, so a plain digit sum is unchanged at 27 either way and would report no problem. The weighted sum, however, becomes $0 + 27 + 0 + 42 + 24 + 0 + 24 + 15 + 2 + 2 = 136$, and $136 - 132 = 4$, so it is not a multiple of 11 and the record is rejected. That is the entire reason check digits are weighted rather than plain: position has to matter, or the commonest human error passes straight through.

Luhn: the number that fails, and the one case it misses: The Luhn check, used on payment and many membership numbers, doubles every second digit counting from the right, subtracts 9 from any doubled result above 9, sums everything, and requires a total ending in zero. Take 4539 1488 0343 6467. The

doubled positions contribute 3, 3, 8, 0, 7, 2, 6 and 8, totalling 37; the undoubled positions contribute 7, 4, 3, 3, 8, 4, 9 and 5, totalling 43. The grand total is 80, which ends in zero, so the number passes. Now transpose the last two digits to 4539 1488 0343 6476. The doubled contributions become 5, 3, 8, 0, 7, 2, 6, 8 = 39 and the undoubled become 6, 4, 3, 3, 8, 4, 9, 5 = 42, giving 81. Not a multiple of ten, so the mistyped number is rejected at the point of entry. Luhn catches every single-digit substitution and almost every adjacent transposition: with exactly one blind spot. Swapping an adjacent 0 and 9 leaves the total unchanged, because doubling 0 gives 0 and doubling 9 gives 18 which reduces to 9, so both digits contribute the same amount whether doubled or not. A form that validates with Luhn will happily accept 90 where you typed 09. Knowing the blind spot is the point: a check digit narrows the space of undetected errors, it never closes it.

Negative marking: 44 right can beat 46 right: Many clerical checking sections score correct minus incorrect, with omissions neutral. Under that rule, guessing between two remaining options has an expected value of $0.5 \times (+1) + 0.5 \times (-1) = 0$, exactly neutral, and guessing blindly among four options has an expected value of $0.25 - 0.75 = -0.5$, clearly negative. Concretely: on a 60-item section you attempt 48, get 44 right and 4 wrong, and score 40. A colleague attempts all 60, gets 46 right and 14 wrong, and scores 32. More correct answers, a lower score. Reverse the scoring rule to plain raw-correct with no penalty and the ordering flips: 46 beats 44 and leaving a blank becomes strictly irrational. Neither strategy is universally right, which is why the instruction screen is not optional reading. Find the sentence that says whether wrong answers are penalised, and if it is absent, assume raw scoring and answer everything.

Reconcile the batch: consistent is not the same as correct: An invoice lists three lines: 12 at $14.25 = 171.00$; 7 at $33.60 = 235.20$; 3 at $128.00 = 384.00$. The stated subtotal is 709.20, VAT at 20 percent is stated as 141.84, and the total is stated as 851.04. Every downstream figure checks out against the one above it: $709.20 \times 0.20 = 141.84$ and $709.20 + 141.84 = 851.04$. Nothing in the tax or total column is inconsistent. But recompute the lines: $171.00 + 235.20 + 384.00 = 790.20$, not 709.20. The subtotal is a transposition of the correct figure, and because every later number was calculated from the wrong subtotal, the document is perfectly self-consistent and perfectly wrong. The true VAT should be 158.04 and the true total 948.24, a shortfall of 97.20. This is the single most valuable habit in the construct: never verify a total against the number printed above it, always recompute it from the underlying items. Internal consistency proves that one person did the arithmetic carefully; it proves nothing about whether they started from the right number.

How to practise this skill

- Drill transposition specifically. Generate pairs where the only difference is two adjacent characters swapped, because that is the family your eye is worst at and the family a naive check will not catch.
- Fix a scan order and never vary it: field by field, top to bottom, and within a long value block by block. Free-roaming comparison feels faster and reliably misses one field per record.
- Learn one check-digit scheme by hand, modulus 11 is the easiest, until you can run it in about twenty seconds. It converts a whole class of 'does this record look right' items from judgement into arithmetic.
- Read the scoring rule before the first item and decide your guessing policy then, not at minute seven when you are behind. Under penalty scoring, an omission is a legitimate answer; under raw scoring it is a wasted mark.
- When you are checking financial or tabular documents, recompute the lowest-level figures first and work upward. Errors introduced at a subtotal propagate perfectly and are invisible from below.
- Log the errors you miss, not the ones you find. A miss log after four sessions usually shows a personal signature (always the middle of long numeric strings, or always the second occurrence of a repeated field), and that signature is what to drill.

Glossary

Transposition error: Two adjacent characters exchanged, as in 496 typed for 469. Length and character content are unchanged, so length checks and plain digit sums both pass it; only positional comparison or a

weighted check digit detects it.

Substitution error: One character replaced by another, as in 08571 for 08871. It is the family most reliably caught by reading a value aloud against the source, and by any check digit scheme.

Check digit: An extra digit computed from the others so that a corrupted record fails an arithmetic test. It detects errors; it never corrects them and never proves the record refers to the right person or item.

Modulus 11: A weighted check scheme in which digits are multiplied by descending position weights and the total must divide by 11. Used in ISBN-10 and several banking and national identifier formats.

Luhn algorithm: A modulus 10 check that doubles alternate digits from the right and requires the sum to end in zero. It catches all single-digit errors and most adjacent transpositions except an adjacent 0 and 9.

Miss and false alarm: A miss is a genuine discrepancy passed as correct; a false alarm is a difference flagged where none exists, usually a formatting variant such as a trailing space or a differently written date. Misses are the costlier of the two in records work, which is why procedures normalise formatting first and then bias toward escalating doubt.

Correction for guessing: A scoring rule that subtracts a fraction or multiple of the wrong answers from the correct ones, making blind guessing negative in expectation and turning omission into a rational choice.

Reconciliation: Recomputing an aggregate from its components rather than accepting the printed figure. It is the only check that catches an error introduced at summary level, where everything below and above it still agrees.

Study this next

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

Where this material comes from

- The ISBN-10 modulus 11 worked example (0-306-40615-2, weighted sum 132) and the Luhn worked example (4539 1488 0343 6467, sum 80) were computed by hand for this lesson and each was re-verified digit by digit, including the corrupted variants.
- Algorithm definitions cross-checked against the public specifications described in the Wikipedia articles 'International Standard Book Number' and 'Luhn algorithm', including the documented 09/90 transposition blind spot. Records, invoices and reference numbers above are invented.
- 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.

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.

Logical reasoning

Conditional logic, syllogisms, ordering, grouping, argument structure, and constraint problems.

Logical reasoning is the ability to take a set of statements (conditionals, quantifiers, ordering and grouping constraints), and work out exactly what they force, what they permit, and what they leave open. It is the backbone of graduate screening batteries, law and policy admissions tests, analyst and investigator selection, and the constraint sections of programming aptitude tests. The same discipline runs through eligibility policy, contract clauses, access-control rules and any specification written as if-then. 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. Translate everyday phrasing (only if, unless, no A is B, none but) into a single arrow of the form antecedent implies consequent, and get the direction right every time.
2. Derive the contrapositive of any conditional and use it to reason backwards from a denied consequent.
3. Name and reject the two invalid conditional moves, affirming the consequent and denying the antecedent, in items designed to make both feel natural.
4. Solve a linear ordering item to a single arrangement by placing the most restrictive block first and eliminating cases exhaustively.
5. Answer a grouping item by identifying who is ruled out, not just who is possible, and distinguish must be true from could be true.
6. Apply quantifier rules correctly: recognise that some means at least one and does not imply not all, and that two overlapping some statements never guarantee a third overlap.

Worked examples and pitfalls

One conditional, four moves, two of them wrong: Take a warehouse rule: if a pallet is flagged by the scanner, then it is inspected before dispatch. Write it as F implies I. Four things can now happen. You learn a pallet was flagged: F is true, so I is true: inspected. That is modus ponens and it is valid. You learn a pallet was not inspected: not-I, so not-F. It was not flagged. That is modus tollens, equally valid, and it is the move most candidates never reach for. Now the two traps. You learn a pallet WAS inspected and conclude it must have been flagged. Invalid: the rule says nothing about why else a pallet might be inspected: random audit, a customer complaint, a damaged corner. That is affirming the consequent. You learn a pallet was NOT flagged and conclude it was not inspected. Invalid for the same reason; that is denying the antecedent. Both feel right because in ordinary conversation people say if when they mean if and only if. Assessment items

exploit exactly that slip, so the fix is mechanical: the instant you meet a conditional, write F implies I on one line and its contrapositive not-I implies not-F underneath. Those two lines are everything the statement gives you. Anything else on the page is a distractor.

Chaining conditionals, then running the chain backwards: Three statements from a release policy. If the build fails, the deploy is blocked. If the deploy is blocked, the release date slips. If the release date slips, the client is notified. Write them: B implies D, D implies S, S implies N. Chaining forwards is easy: a failed build guarantees a notified client. The item almost never asks that. It says: the client was not notified. What follows? Take contrapositives and chain them the other way: not-N implies not-S, not-S implies not-D, not-D implies not-B. So the build did not fail, the deploy was not blocked, and the date did not slip. All three follow. Now the near-miss version, and the one candidates get wrong: the client WAS notified. What follows? Nothing about the build. The chain only runs one way, and a client can be notified for reasons the three statements never mention. Answer: none of the above must be true. A useful habit for chain items is to draw the arrows on one line, B implies D implies S implies N, and remember that you can travel left-to-right when you are given a truth and right-to-left when you are given a falsehood, never the reverse.

The four phrases that flip the arrow: Most lost marks in conditional items come from translation, not from reasoning. Only if introduces the consequent: you may board only if you hold a boarding pass means board implies pass. It does NOT mean a pass gets you on the plane. The pass is necessary, not sufficient, and you still need the gate to be open and your name off the no-fly list. Unless is read as if not: unless you check in you cannot board is not-check-in implies not-board, whose contrapositive is board implies check-in, again the check-in is necessary. None but staff may enter is enter implies staff. No temporary badge grants server-room access is temporary badge implies not server-room access. Compare all four with a plain sufficient condition: swiping a manager card opens the door is card implies open, and here the card really is enough. The discipline is to ask one question of every clause. Is this thing the guarantee or the requirement? A guarantee sits at the tail of the arrow, a requirement sits at the head. Get that one decision right and the rest of the item is bookkeeping.

An ordering item, solved to a single arrangement: Five people present in five consecutive slots, one each: Aisha, Ben, Chen, Dana, Eli. Constraints: (1) Ben presents in the slot immediately before Dana. (2) Aisha is neither first nor last. (3) Chen presents at some point before Ben. (4) Eli is not immediately before or immediately after Dana. (5) Dana is not last. Start with the most restrictive item, the Ben-Dana block, and test each placement. Slots 1-2 is impossible: Chen must come before Ben and there is no slot before 1. Slots 4-5 is impossible because Dana would be last, breaking (5). Slots 3-4: Chen takes 1 or 2, and Aisha must take 2 because she cannot be first or last, so Chen is 1 and Eli is forced into 5, but 5 is immediately after Dana in slot 4, breaking (4). Dead. Slots 2-3: Chen must be 1. Aisha and Eli take 4 and 5, and since Aisha cannot be last, Aisha is 4 and Eli is 5. Check (4): Dana is in 3, the neighbouring slots are 2 and 4, and Eli is in 5: fine. The order is Chen, Ben, Dana, Aisha, Eli, and it is the only one. Two habits made that quick. First, place the block before the singletons; a two-slot block only has four possible positions in a five-slot line, so it does most of the elimination for you. Second, when a case dies, write down which constraint killed it. Later questions in the same set often ask what happens if one rule is dropped, and your dead-case notes answer them instantly.

A grouping item where the real answer is who is excluded: Exactly three of six people are picked: Farah, Gus, Hana, Ivo, Jo, Kit. Rules: (1) If Farah is picked, Gus is not. (2) Hana is picked only if Ivo is. (3) At least one of Gus and Jo is picked. (4) Kit and Ivo are never both picked. The question: if Hana is picked, which of the following must be true? Work it. Rule (2) is Hana implies Ivo, so Ivo is in. That is two of the three seats. Rule (4) then puts Kit out. The third seat goes to Farah, Gus or Jo. Suppose it went to Farah: then neither Gus nor Jo is picked, breaking (3). So the third seat is Gus or Jo, and both of those complete a legal team. Check Gus: rule (1) is vacuous because Farah is out, (2) holds, (3) holds, (4) holds. Check Jo: same. So the answer is not who joins Hana; that is genuinely open. The answer is that Farah is NOT picked, and it must be true in every legal arrangement. Grouping items are built around that distinction. Could be true means one

arrangement exists; must be true means no counter-arrangement exists. The fastest way to kill a must-be-true option is to build a single legal team that violates it, which is why sketching two complete valid teams before reading the options is usually faster than reasoning about the options one at a time.

Some, all and none: the overlap nobody gave you: Two premises: all compliance officers have completed the ethics module, and some people who completed the ethics module work in the Leeds office. Does it follow that some compliance officers work in Leeds? No, and the diagram shows why. Draw a big circle for ethics-module completers. Compliance officers sit entirely inside it. Leeds staff overlap it somewhere, but nothing places that overlap on top of the compliance officers, so a world where every Leeds completer is a lab technician satisfies both premises and refutes the conclusion. Two overlapping particular statements never chain. Now three rules worth memorising. First, some means at least one and is silent about the rest, so the premise some invoices are overdue does not rule out all of them being overdue, though it does not establish that either. An option reading all invoices are overdue is therefore could-be-true rather than must-be-true, and a candidate who takes some to mean some but not all will wrongly mark it impossible. Second, some does convert: if some completers are Leeds staff then some Leeds staff are completers, and that is valid. All does not convert. All compliance officers are completers tells you nothing about most completers. Third, a valid chain needs a universal: all A are B plus no B are C gives no A are C, and that one is airtight. When an item mixes quantifiers, sketch three circles before you read a single answer option; the arrangement that satisfies the premises while breaking the conclusion is usually visible in about ten seconds.

How to practise this skill

- Notate before you reason. Every conditional gets two lines on your paper, the statement and its contrapositive, written in symbols. Items that look like paragraphs collapse to four or five arrows, and the arrows are what the question is actually about.
- Underline only if, unless, none but and no as you read. Those four phrasings cause more errors than every inference rule combined, and they are the cheapest thing on the page to get right.
- On ordering sets, draw a numbered row of slots and place the largest fixed block first. Keep a note of which constraint killed each dead case; follow-up questions in the same set reuse them.
- On must-be-true options, attack rather than verify. One complete legal arrangement that breaks the option kills it outright, and building one is usually faster than proving the option holds everywhere.
- Train the invalid moves deliberately. Write out an affirming-the-consequent item and a denying-the-antecedent item of your own each session until the shape of them is instantly recognisable under time pressure.
- Work untimed until you can state which rule licensed each step, then add the clock. Attempts are stored on this device only, so a slow, fully-narrated first pass through a constraint set costs nothing but your own time.

Glossary

Antecedent and consequent: In a conditional if P then Q, P is the antecedent and Q is the consequent. Getting the two the right way round during translation is where most conditional items are won or lost.

Contrapositive: For P implies Q, the statement not-Q implies not-P. It is always logically equivalent to the original, which is what lets you reason backwards from a denied consequent.

Modus ponens: The valid inference from P implies Q together with P to the conclusion Q. The most common inference in assessment items, and the safest.

Modus tollens: The valid inference from P implies Q together with not-Q to the conclusion not-P. Under-used by candidates, and the move most backwards-chaining items are built around.

Affirming the consequent: The invalid move from P implies Q together with Q to the conclusion P. It is tempting because ordinary speech often means if and only if when it says if.

Necessary condition: Something that must hold for an outcome to occur, but does not by itself bring it about. Signalled by only if, unless, none but, and required for.

Sufficient condition: Something whose presence guarantees the outcome. Signalled by a plain if, whenever, or any time that. A condition can be necessary, sufficient, both, or neither.

Must be true versus could be true: Must be true holds in every arrangement the constraints permit; could be true holds in at least one. Nearly every grouping and ordering answer option is one of these two, and mistaking which is being asked costs the mark.

Study this next

- Logical reasoning skill hub: <https://learn.novusstreamsolutions.com/aptitude/skills/logical-reasoning>
- Logical reasoning practice bank: <https://learn.novusstreamsolutions.com/cognitive-skills/logical-reasoning>
- Deductive reasoning lesson: <https://learn.novusstreamsolutions.com/aptitude/skills/deductive-reasoning/lesson>
- Programming logic suite: <https://learn.novusstreamsolutions.com/aptitude/tests/programming-logic>
- Law, compliance and justice careers: <https://learn.novusstreamsolutions.com/careers/families/law-compliance-and-justice>
- Logical reasoning lesson on Novus Learn: <https://learn.novusstreamsolutions.com/aptitude/skills/logical-reasoning/lesson>

Where this material comes from

- Worked items written for Novus Learn from standard introductory logic: conditional translation, the contrapositive, categorical syllogisms, and linear ordering under constraints. Every puzzle above was solved and checked for a unique answer before publication.
- Terminology checked against the public Wikipedia articles 'Modus ponens', 'Modus tollens', 'Contraposition', 'Affirming the consequent' and 'Syllogism'. Definitions only; all items are 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 reading comprehension: <https://learn.novusstreamsolutions.com/aptitude/skills/reading-comprehension/lesson>
- Practice critical thinking: <https://learn.novusstreamsolutions.com/aptitude/skills/critical-thinking/lesson>
- Practice error detection: <https://learn.novusstreamsolutions.com/aptitude/skills/error-detection/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/contract-review/study-outline>