

Policy and program analysis: study guide

Government and public policy · suite apt-003-policy-and-program-analysis · generated
2026-09-15T13:31:21.817Z

Detail	Value
Title	Policy and program analysis: study guide
Generated	2026-09-15T13:31:21.817Z
Fixture/version	apt-003-policy-and-program-analysis
Sector	Government and public policy
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 Policy and program analysis

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-003-policy-and-program-analysis&mode=guided>
- Start mini-test:
<https://learn.novusstreamsolutions.com/aptitude/practice/new?bank=apt-003-policy-and-program-analysis&mode=mini>
- Start full simulation:
<https://learn.novusstreamsolutions.com/aptitude/practice/new?bank=apt-003-policy-and-program-analysis&mode=full>

Skills covered, in full

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

Data interpretation

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

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

What you should be able to do after this lesson:

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

Worked examples and pitfalls

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

One pair of rows, three correct increases: A complaints table: 2023, 120,000 orders, complaint rate 4.0 percent; 2024, 150,000 orders, complaint rate 5.0 percent. Three defensible answers to 'how much did complaints increase?'. The rate rose by 1.0 percentage point. The rate rose by $(5.0 - 4.0) / 4.0 = 25$ percent in relative terms. And the count of complaints rose from $0.04 \times 120,000 = 4,800$ to $0.05 \times 150,000 = 7,500$, which is $(7,500 - 4,800) / 4,800 = 56.25$ percent. All three are arithmetically right; only one answers the question in front of you. The pattern to internalise is that a rate and a count move together only when the denominator is fixed, and here it is not. Order volume grew 25 percent as well. If the question is about

customer experience, the rate is the honest figure; if it is about how many complaint handlers to hire, the count is. Test items usually ask for the one you would not have chosen.

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

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

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

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

How to practise this skill

- Read the title, the units line, the row and column headers and any footnote before you look at a single value. Roughly the first fifteen seconds of an item should contain no arithmetic at all, and that fifteen seconds is what prevents the thousand-fold and percentage-point errors.
- For every item, write down which of the three quantities is wanted (share of total, relative change, or change in percentage points), before computing. Most wrong answers on this construct are correct arithmetic applied to the wrong quantity.
- Whenever two groups differ in size, ask what the denominator should be. If a question compares sites,

teams, countries or periods of unequal length and you have not divided by something, you are almost certainly answering the wrong question.

- Practise the cannot-say items separately and force yourself to name the missing variable each time. Candidates who train only on computational items reliably over-answer inference statements under time pressure.
- Do not redraw or re-scale charts in your head. Locate the gridline values or the data labels, and if neither exists, interpolate between two labelled gridlines and state the bound rather than guessing a point value.
- Time yourself per item rather than per section. Data interpretation sets share a stimulus, so the first item costs the reading time and the rest should be fast; if item four takes as long as item one, you did not build a mental map of the table.

Glossary

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

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

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

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

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

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

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

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

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

Study this next

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

Where this material comes from

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

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

Critical thinking

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

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

What you should be able to do after this lesson:

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

Worked examples and pitfalls

The negation test finds assumptions; nothing else does: Argument: 'The council should fit bin sensors across the district. In the trial depot they cut collection trips by a fifth.' Candidate assumption A: 'The trial depot's waste pattern is broadly representative of the rest of the district.' Negate it - suppose the trial depot is nothing like the rest of the district - and the argument collapses, because a result that does not transfer supports nothing about the district. So A is a required assumption. Candidate assumption B: 'Bin sensors are the cheapest available technology for this purpose.' Negate it - suppose they are the most expensive - and the argument still stands as given, since it argued from a reduction in trips, not from cost. B is not required; it would strengthen a cost-based argument that was never made. Candidate C: 'Fewer collection trips reduce total operating cost.' Negate it and the recommendation loses its point, so C is required too, even though the argument never mentions cost. That last case is the one candidates miss, because the required assumption

is what bridges the evidence to the recommendation, and bridging assumptions are by definition unstated. The failure mode across this whole item type is picking statements that support the conclusion instead of statements the argument cannot survive without.

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

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

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

Interpretation: what a survey figure does and does not license: Data: all 1,200 employees of a firm were surveyed. 62 percent said they would accept a four-day week at 90 percent pay. Of those who said yes, 71 percent were under 35. Conclusion one: 'A majority of this firm's employees would accept the trade.' This follows - 62 percent of a complete census of the firm is a majority, and the conclusion is properly limited to this firm. Conclusion two: 'A majority of the firm's under-35 employees would accept the trade.' This does not follow, and the arithmetic shows why. Sixty-two percent of 1,200 is 744 accepters, and 71 percent of 744 is about 528 accepters aged under 35. That is the composition of the accepters, not the acceptance rate within

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

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

How to practise this skill

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

Glossary

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

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

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

Affirming the consequent: The invalid pattern 'if P then Q; Q; therefore P'. It is the most common deductive distractor because it differs from the valid modus tollens form by only the position of a negation.

Base rate: How common something is in the population before any test or evidence is applied. Ignoring it

makes accurate-sounding tests appear far more informative than they are when the condition is rare.

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

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

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

Study this next

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

Where this material comes from

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

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

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.

Case analysis

Synthesizing multiple documents and making a reasoned recommendation.

Build a recommendation from several documents without blurring facts, interpretations, and assumptions. The core method is to frame the decision, organize evidence, test alternatives, and state conditions or gaps.

What you should be able to do after this lesson:

1. Translate a broad case prompt into a decision question, criteria, constraints, and required output.
2. Create an evidence table that separates supported findings from inference and missing information.
3. Write a recommendation that compares alternatives, acknowledges trade-offs, and names the next validation step.

Worked examples and pitfalls

Worked scenario: two service proposals: Proposal A is cheaper and can start now; Proposal B costs more but has stronger reliability evidence. The case does not provide peak-demand data. State the decision criteria first (cost, launch timing, reliability, and demand risk), then compare only supported facts. A defensible recommendation might select a limited A pilot with a reliability threshold and request peak-demand evidence before wider rollout, rather than claiming either proposal is universally best.

Evidence matrix: Use columns for claim, supporting document, strength, contradiction, and gap. When sources conflict, report the conflict and explain how it affects confidence. Repeating every document is summary; connecting evidence to the decision criteria is analysis.

How to practise this skill

- Begin with the decision and criteria so interesting but irrelevant details do not take over.
- Use calibrated language: supported, suggests, uncertain, and conditional communicate different evidence strength.
- End with an owner, next step, and condition for revisiting the recommendation.

Glossary

Finding: A conclusion directly supported by the evidence supplied in the case.

Inference: A reasoned interpretation that goes beyond a directly stated fact and therefore needs justification.

Decision criterion: A standard used to compare options, such as cost, safety, feasibility, equity, or timing.

Study this next

- Case study: <https://learn.novusstreamsolutions.com/search?q=Case%20study>
- Evidence: <https://learn.novusstreamsolutions.com/search?q=Evidence>
- Recommendation: <https://learn.novusstreamsolutions.com/search?q=Recommendation>
- Case analysis lesson on Novus Learn:
<https://learn.novusstreamsolutions.com/aptitude/skills/case-analysis/lesson>

Where this material comes from

- Novus Learn original case-study, policy-analysis, investigation, and evidence suite scenarios.
- Novus educational framework: frame the decision, map evidence, compare alternatives, qualify confidence, and recommend.

Situational judgment

Evaluating workplace responses against role-relevant principles.

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

What you should be able to do after this lesson:

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

Worked examples and pitfalls

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

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

How to practise this skill

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

Glossary

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

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

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

Study this next

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

Where this material comes from

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

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

Recommended preparation

- Practice data interpretation:
<https://learn.novusstreamsolutions.com/aptitude/skills/data-interpretation/lesson>
- Build critical thinking: <https://learn.novusstreamsolutions.com/aptitude/skills/critical-thinking/lesson>
- Learn case analysis: <https://learn.novusstreamsolutions.com/aptitude/skills/case-analysis/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/policy-and-program-analysis/study-outline>