

Crime and intelligence analysis: study guide

Police, justice, corrections, border, and security · suite apt-049-crime-and-intelligence-analysis · generated 2026-09-15T15:27:45.004Z

Detail	Value
Title	Crime and intelligence analysis: study guide
Generated	2026-09-15T15:27:45.004Z
Fixture/version	apt-049-crime-and-intelligence-analysis
Sector	Police, justice, corrections, border, and security
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 Crime and intelligence 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-049-crime-and-intelligence-analysis&mode=guided>
- Start mini-test:
<https://learn.novusstreamsolutions.com/aptitude/practice/new?bank=apt-049-crime-and-intelligence-analysis&mode=mini>
- Start full simulation:
<https://learn.novusstreamsolutions.com/aptitude/practice/new?bank=apt-049-crime-and-intelligence-analysis&mode=full>

Skills covered, in full

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

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.

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.

Inductive reasoning

Identifying general rules from examples, sequences, and observations.

Inductive reasoning is the ability to look at a handful of examples (a number series, a set of observations, a table of past decisions), and propose the rule that generated them, while staying honest about the fact that finished evidence never fully proves a rule. It carries most of the weight in graduate ability batteries, analyst and intelligence screening, and college placement testing, and it is what a technician does when three failures on the same line suggest a pattern rather than three coincidences. The construct rewards two opposite habits at once: generate candidate rules quickly, then try hardest to break the one you like. Everything you practise here is device-local unless you export it.

What you should be able to do after this lesson:

1. Compute first and second differences on a numeric series and use them to distinguish arithmetic, quadratic and multiplicative rules.
2. Detect interleaved series, where alternate terms follow two separate rules, before concluding a sequence has no pattern.
3. Convert letter sequences to alphabet positions, apply the same difference machinery, and wrap correctly past position 26.
4. Explain why a finite sequence is under-determined, and choose between competing rules by testing each against every given term rather than the first two.

5. Induce a decision rule from a small table of labelled cases, and identify the single row that rules out the obvious over-general hypothesis.
6. Design the test that could disconfirm your hypothesis, and recognise why confirming cases are worth far less than disconfirming ones.

Worked examples and pitfalls

Differences first, ratios second: Four series, each solved by the same opening move. Series one: 2, 6, 12, 20, 30. First differences are 4, 6, 8, 10: arithmetic, rising by 2, so the next difference is 12 and the next term is 42. The closed form is n times n -plus-one: $1 \times 2, 2 \times 3, 3 \times 4, 4 \times 5, 5 \times 6, 6 \times 7 = 42$, which confirms it. Series two: 7, 10, 16, 28, 52. First differences are 3, 6, 12, 24. Those are doubling, so the next difference is 48 and the next term is 100. Equivalently each term is double the previous minus 4: $7 \times 2 - 4 = 10, 10 \times 2 - 4 = 16, 16 \times 2 - 4 = 28, 28 \times 2 - 4 = 52, 52 \times 2 - 4 = 100$. Series three: 3, 4, 7, 11, 18, 29. Nothing simple in the differences, so try summing neighbours: $3 + 4 = 7, 4 + 7 = 11, 7 + 11 = 18, 11 + 18 = 29$, and the next term is $18 + 29 = 47$. Series four: 1, 4, 9, 61, 52. That looks broken until you notice 61 and 52 are 16 and 25 with their digits reversed; the underlying series is the squares 1, 4, 9, 16, 25, 36, written back-to-front where reversing changes anything, so the next term is 63. The order of attack that solves most series is fixed: first differences, then second differences, then ratios, then sums of neighbours, then digit manipulation. Run it in that order and you stop staring.

Two rules fit; the fourth term decides: You are shown 2, 4, 8 and asked for the next term. Doubling gives 16. But first differences are 2 and 4, and if those differences are themselves rising by 2 the next difference is 6 and the next term is 14. Both rules fit every term you were given, so the sequence is genuinely under-determined and neither answer is more correct than the other from the data alone. This is not a trick; it is the central limitation of induction, and assessment writers manage it by giving you enough terms to separate the candidates. Add one term. If the series is 2, 4, 8, 14 then the differences are 2, 4, 6 and the next difference is 8, giving 22, doubling is dead, because doubling would have produced 16 at position four. If the series is 2, 4, 8, 16 the difference rule is dead instead, because it predicted 14. So the practical rule is: never commit on two terms, always test your candidate rule against the LAST given term as well as the first, and if two rules survive all given terms, look at the answer options, exactly one of them will normally correspond to a rule that fits, and that is legitimate evidence about what the writer intended.

Letters are numbers wearing a costume: The sequence C, F, J, O, U. Convert to alphabet positions: C is 3, F is 6, J is 10, O is 15, U is 21. First differences are 3, 4, 5, 6, rising by one, so the next gap is 7 and the next position is 28. The alphabet has only 26 letters, so wrap: 28 minus 26 is 2, which is B. Answer B. Three things go wrong on letter items. Candidates count gaps by reciting the alphabet on their fingers and drop a letter, which is why writing the numbers down beats counting in your head every time. Candidates forget to wrap, and answer with a position number rather than a letter. And candidates mishandle a series that runs backwards past A. Position 1 minus 3 is minus 2, which wraps to 26 minus 2, that is 24, the letter X. A related family uses letter pairs where the two letters move at different rates, for example AZ, CX, EV, GT: the first letters are 1, 3, 5, 7 going up by two and the second are 26, 24, 22, 20 going down by two, so the next pair is I and R, that is IR. Split the pair, treat each stream separately, and the item becomes two easy arithmetic series instead of one impossible one.

Interleaving: two clocks in one series: The series 5, 8, 6, 11, 7, 14, 8. First differences are 3, minus 2, 5, minus 4, 7, minus 6: alternating sign, no obvious progression, and this is the point where candidates guess. Split it instead. Terms in the odd positions are 5, 6, 7, 8, rising by one. Terms in the even positions are 8, 11, 14, rising by three. The next term sits in an even position, so it continues the second stream: 14 plus 3 is 17. The diagnostic that should trigger the split is alternating signs in the first differences, or a series that oscillates while drifting. Interleaving also appears with three streams, and with one stream constant: 4, 9, 4, 16, 4, 25 hides the squares 9, 16, 25 among repeated 4s, so the next term is 4 and the one after is 36. One caution worth carrying: a series that has been split should be checked back against the original by writing out

your predicted term in place and reading the whole thing again. If the reassembled series looks stranger than the one you started with, the split was probably wrong.

Inducing a rule from labelled cases: Five support tickets, each with a customer tier, a first-response time, and whether it was escalated. Ticket 1: Enterprise, 6 hours, escalated. Ticket 2: Enterprise, 2 hours, not escalated. Ticket 3: Standard, 9 hours, not escalated. Ticket 4: Enterprise, 5 hours, escalated. Ticket 5: Standard, 1 hour, not escalated. Three hypotheses are worth writing down. Hypothesis A: escalation happens when the response is slow. Ticket 3 kills it. A nine-hour Standard ticket was not escalated. Hypothesis B: escalation happens for Enterprise customers. Ticket 2 kills it. A fast Enterprise ticket was not escalated. Hypothesis C: escalation requires BOTH Enterprise tier and a response over four hours. It fits all five rows, and it is the only conjunction that does. Notice which rows did the work: the two that fit one hypothesis but not the label are worth more than the three that agree with everything. Now the honest limit. Does an Enterprise ticket answered in exactly four hours escalate? Hypothesis C as written says no, but the data contains no case between two and five hours, so the true threshold could be anywhere in that window. Induction from five rows gives you a rule and a region of ignorance, and naming the region is part of the answer.

Look for the case that would break you: Four cards lie on a table showing A, K, 4 and 7. Each card has a letter on one side and a number on the other. The rule to test: if a card has a vowel on one face, it has an even number on the other. Which cards must you turn? Most people say A, and many add 4. A is right: a vowel with an odd number on the back refutes the rule directly. The 4 is useless: the rule says nothing about what must be behind an even number, so a consonant there breaks nothing and a vowel there merely agrees. The card people miss is the 7. Turn it, find a vowel, and the rule is dead. So the answer is A and 7, the two cards that could produce a violation. This is the Wason selection task, a well-known public result in the psychology of reasoning, and the reason it belongs in an inductive lesson is that it isolates the habit the construct actually rewards: people search for confirming evidence by default and have to be trained to search for disconfirming evidence. Carry it into series items as a check. Once you have a candidate rule, do not run it forward on the terms it was built from; run it on the term you have not used yet, and treat a single mismatch as fatal rather than as noise.

How to practise this skill

- Write the first differences under every numeric series before you think about it at all. The mechanical step surfaces arithmetic and quadratic rules for free and tells you within seconds whether to move on to ratios.
- Fix an order of attack and follow it: first differences, second differences, ratios, sums of adjacent terms, then digit tricks. Most lost time in this construct is spent re-staring at a series rather than working through the list.
- Convert letters to numbers on paper the moment you see a letter series, and write the alphabet with positions at the top of your rough sheet once at the start of a session rather than counting on your fingers per item.
- Treat alternating signs in the first differences as a direct instruction to split the series into odd and even positions. The interleaved family is large and is where candidates most often decide, wrongly, that there is no pattern.
- Validate a candidate rule against the last given term, not the first two. A rule that explains the opening of a series and misses its final term is the standard wrong answer, and the answer options are usually built to reward it.
- When you induce a rule from a table of cases, name the row that eliminated the simpler hypothesis and name the gap the data leaves open. Both are recorded on this device only, so build the habit here where over-claiming is free.

Glossary

Inductive inference: Reasoning from particular observations to a general rule. The conclusion is supported but never guaranteed, which is the essential difference from deduction.

First difference: The gap between consecutive terms of a series. A constant first difference means an arithmetic progression; a constant second difference means a quadratic rule.

Geometric progression: A series where each term is the previous term multiplied by a fixed ratio. Detected by dividing neighbours rather than subtracting them.

Interleaved series: A single printed sequence containing two or more independent series in alternate positions. Signalled by first differences that alternate in sign.

Under-determination: The condition in which several different rules fit every term you were given, so the data alone cannot single one out. It is the reason short series need extra terms or answer options to be decidable.

Disconfirming instance: A case that a hypothesis predicts should not exist. One of these refutes a rule outright, whereas any number of agreeing cases only fails to refute it.

Confirmation bias: The tendency to seek evidence that agrees with a hypothesis already held. In series and rule-induction items it shows up as checking a rule against the terms that suggested it.

Wrap-around: Continuing a letter sequence past Z back to A, or before A back to Z, by adding or subtracting 26 from the alphabet position. The routine source of near-miss answers on letter items.

Study this next

- Inductive reasoning skill hub: <https://learn.novusstreamsolutions.com/aptitude/skills/inductive-reasoning>
- Inductive reasoning practice bank:
<https://learn.novusstreamsolutions.com/cognitive-skills/inductive-reasoning>
- Abstract reasoning lesson:
<https://learn.novusstreamsolutions.com/aptitude/skills/abstract-reasoning/lesson>
- Ability-test style familiarization suite:
<https://learn.novusstreamsolutions.com/aptitude/tests/shl-style-ability-test-familiarization>
- Crime and intelligence analysis suite:
<https://learn.novusstreamsolutions.com/aptitude/tests/crime-and-intelligence-analysis>
- Inductive reasoning lesson on Novus Learn:
<https://learn.novusstreamsolutions.com/aptitude/skills/inductive-reasoning/lesson>

Where this material comes from

- Worked series written for Novus Learn from standard sequence analysis: first and second differences, geometric ratios, neighbour sums, alphabet-position arithmetic, and interleaved streams. Every series above was extended and checked term by term.
- The four-card selection problem is the Wason selection task, first published by P. C. Wason in 1968 and widely reproduced in the public psychology-of-reasoning literature; the description above is a plain restatement of the standard version, not an item from any commercial test.
- Terminology checked against the public Wikipedia articles 'Inductive reasoning', 'Arithmetic progression', 'Geometric progression' and 'Wason selection task'. Definitions only.
- 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.

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.

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.

Recommended preparation

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

Answer keys, scoring and privacy

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

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

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

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

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

Source: <https://learn.novusstreamsolutions.com/aptitude/tests/crime-and-intelligence-analysis/study-outline>