

Compliance and anti-money-laundering: study guide

Finance, accounting, banking, and insurance · suite apt-244-compliance-and-anti-money-laundering · generated 2026-09-15T15:42:57.480Z

Detail	Value
Title	Compliance and anti-money-laundering: study guide
Generated	2026-09-15T15:42:57.480Z
Fixture/version	apt-244-compliance-and-anti-money-laundering
Sector	Finance, accounting, banking, and insurance
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 Compliance and anti-money-laundering

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-244-compliance-and-anti-money-laundering&mode=guided>
- Start mini-test:
<https://learn.novusstreamsolutions.com/aptitude/practice/new?bank=apt-244-compliance-and-anti-money-laundering&mode=mini>
- Start full simulation:
<https://learn.novusstreamsolutions.com/aptitude/practice/new?bank=apt-244-compliance-and-anti-money-laundering&mode=full>

Skills covered, in full

This suite draws on 6 skill constructs. Each one below carries its complete lesson.

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.

Attention and concentration

Sustained focus, selective attention, and accuracy under time pressure.

Attention and concentration is the skill of still noticing on minute forty of a checking block as reliably as you did on minute two, and of pointing the noticing at the right thing when two things compete. It has three distinguishable parts - selective attention, sustained attention or vigilance, and divided attention - and assessments load them differently: cancellation and checking tasks load vigilance, conflict tasks load selection, and dispatch-style monitoring loads division. It is the construct that decides whether a records clerk, a dispatcher, an air-side controller or a quality inspector catches the one wrong digit in a shift. Practice here is device-local, with no account and nothing uploaded.

What you should be able to do after this lesson:

1. Count occurrences of a target in a dense string without double-counting or losing your place, using a fixed scan path rather than free looking.
2. Separate hits, misses and false alarms in your own results and work out whether you are set too eager or too cautious for the scoring rule in force.
3. Predict where in a long block your error rate will rise, and schedule micro-resets against that curve instead of pushing through it.
4. Recognise a transposition error, the class that survives a same-characters gist check and is therefore

missed most often.

5. Explain why an automatic process such as reading interferes with a controlled one, and use that to anticipate which distractors will actually cost you time.
6. Run a two-stream monitoring task by alternating on a fixed cadence rather than attempting genuine simultaneity.

Worked examples and pitfalls

A cancellation count you can actually check: Count every 7 in this row: 4 7 1 7 7 3 9 7 2 8 7 5 7 6 0 7.

Working left to right and tapping once per hit: hits at the second, fourth, fifth, eighth, eleventh, thirteenth and sixteenth positions. That is seven sevens. Two error modes produce nearly all the wrong answers. The first is the adjacent pair - the 7 7 at positions four and five gets counted once, because the eye takes a repeated character as one perceptual object. The second is losing the place after the 9, where the visually similar 9 and 7 force a moment of re-checking and the scan restarts a character early, producing an over-count of eight. The defence for both is a fixed scan path with a physical anchor: a fingertip or cursor moving one character at a time, never jumping back. Re-scanning to check is the thing that creates the double-count, so if you must verify, verify by counting a second time from the RIGHT-hand end and comparing totals, never by re-reading part of the row.

Hits, misses and false alarms: the eager candidate loses: A checking batch contains 300 records, 40 of which are genuinely faulty. Candidate A flags 46 records and 36 of them are truly faulty. Hits 36, misses 40 minus 36 equals 4, false alarms 46 minus 36 equals 10. Candidate B flags 38 and 34 are truly faulty: hits 34, misses 6, false alarms 4. On hit rate alone A wins, 36 out of 40 which is 90 percent against B's 34 out of 40 which is 85 percent. Now apply the scoring rule that most checking tasks actually use, where a false alarm cancels a hit: A scores 36 minus 10 equals 26, B scores 34 minus 4 equals 30. B wins by four despite catching two fewer faults. This is why 'flag anything that looks odd' is bad advice on a scored checking task, and why the first thing to read on a checking item is whether wrong flags are penalised. Your response criterion - how much evidence you demand before flagging - is adjustable, and it should be set from the scoring rule, not from your temperament.

The transposition that passes every gist check: Compare these two lines and decide whether they match. Invoice 4820-7391-06. Invoice 4820-7931-06. They do not: the middle group reads 7391 in the first and 7931 in the second, with the 3 and the 9 swapped. Transpositions are the most-missed error class in record checking for a structural reason - the character SET is identical, the length is identical, the first and last characters of the group are identical, so every fast check the visual system runs comes back clean. Substitutions and omissions change the character inventory and get caught; transpositions do not. Two habits raise the catch rate. Read digits in fixed groups of two rather than as a whole number, so 73-91 against 79-31 becomes a mismatch at the first group instead of a subtle difference somewhere in a four-digit blur. And check groups in a deliberately non-natural order - last group, first group, middle group - because reading left to right lets the confirmation you built at the start carry you through the middle, which is exactly where the swap is usually planted.

Conflict: why reading fights you: The classic demonstration is Stroop's, published in 1935: the word RED printed in blue ink, with the instruction to name the ink colour. Naming takes measurably longer, and errors go up, because reading a familiar word is automatic and cannot be switched off, so the automatic response has to be suppressed before the controlled one can be produced. The same conflict has a numeric version you can test on yourself in a second: how many characters are in the string 4 4 4? The answer is three, and the digit 4 pulls at you the entire way. In an assessment this appears wherever the salient feature and the asked-for feature come apart - a chart where the tallest bar is not the answer to the question, a form where the highlighted field is not the one being verified, a row where the bold total is not what the stem requested. The practical move is to name the target feature out loud before you look - 'ink colour', 'character count', 'the value for March' - because pre-loading the target biases selection before the automatic reading response

gets a chance to win.

Where the errors actually appear in a 45-minute block: Errors in a long checking block are not spread evenly. Mackworth's 1948 clock-watching study established the pattern that gives the effect its name: detection declines over a prolonged watch, with the sharpest deterioration early rather than at the very end. Practically, on a 45-minute self-timed checking block, expect your per-minute error rate in minutes 20 to 45 to run visibly above minutes 1 to 20 even though nothing about the material changed, and expect the subjective sense of effort to lag the actual decline, so it will not feel like you are getting worse. Two things work against it. Break the block into three fifteen-minute segments with a ten-second reset between them - look away, unfocus, breathe out - which costs thirty seconds of a 45-minute block, about one percent of the time, and buys back more than that in caught errors. And score your practice by segment rather than as one number, because a single overall accuracy figure hides exactly the information you need: whether your problem is skill, which shows as flat error rate, or endurance, which shows as a rising one.

Two streams: alternate on a cadence, do not try to merge: A dispatch-style monitoring task: keep a running total of the numbers announced on channel one while watching channel two for the code word AMBER. Genuine simultaneity is not available - the two tasks compete for the same control resource - so the choice is not whether to alternate but whether to alternate deliberately or accidentally. Deliberate looks like this: fix the arithmetic to a rhythm, updating the total only at each announcement and holding a single number between updates, which frees the gaps for channel two. If the total is 34 and the next announcement is 7, you spend under a second reaching 41 and then you are free again. Accidental looks like re-deriving the running total from the beginning because you did not trust it, which locks up the whole window and is when AMBER goes past unheard. The measurable failure of divided attention is almost never a failure to hear the target; it is a failure to have any spare capacity at the moment it arrived. The related phenomenon worth knowing is inattention blindness, illustrated by Simons and Chabris in 1999: an unexpected and perfectly visible event is missed entirely when attention is committed to a counting task.

How to practise this skill

- Fix your scan path before you start, and never re-scan backwards to double-check. Backward re-scanning is what produces both double-counts and lost places; if you need to verify a count, run it again from the other end and compare.
- Read the scoring rule before the first item. Whether false alarms are penalised changes the correct strategy completely, and a criterion set for the wrong rule costs more marks than slow reading does.
- Check numeric strings in fixed groups of two or three and in a deliberately shuffled group order. Transpositions are the errors that survive whole-string comparison, and they are the ones planted on purpose.
- Time your practice in segments and log accuracy per segment, not per session. A flat error rate means you need speed; a rising one means you need endurance and breaks, and the two problems have opposite fixes.
- Name the target feature out loud before each conflict item. Pre-loading the relevant dimension is the cheapest available defence against the salient-but-wrong answer.
- Practise with an actual distractor present rather than in silence. A test hall has movement, coughing and page-turning, and attention practice in perfect quiet trains a condition you will not meet.

Glossary

Selective attention: Prioritising one source or feature while suppressing competing ones. It is what fails when the visually salient element of a display captures the response instead of the requested one.

Sustained attention (vigilance): Maintaining detection performance over a long, low-event period. It is the capacity that checking, monitoring and inspection tasks are really sampling.

Vigilance decrement: The decline in detection performance over a prolonged watch, typically steepest early in the block. Mackworth's 1948 clock test is the standard public reference.

Divided attention: Handling two streams that compete for control. Performance is best modelled as fast alternation with a switch cost, not as genuine parallelism.

Hit, miss, false alarm, correct rejection: The four outcomes of any detection decision: flagging a real fault, missing one, flagging a clean record, and correctly leaving a clean record alone. Any honest accuracy claim needs at least the first three.

Response criterion: How much evidence you require before flagging. Shifting it trades misses against false alarms without changing your underlying sensitivity, which is why it must be set from the scoring rule.

Stroop interference: The slowing that occurs when an automatic response, such as reading a word, conflicts with the requested response, such as naming its ink colour. Named for Stroop's 1935 experiments.

Inattentional blindness: Failing to notice a fully visible, unexpected event because attention is committed elsewhere - the effect Simons and Chabris demonstrated in 1999 with a counting task.

Study this next

- Attention and concentration skill hub:
<https://learn.novusstreamsolutions.com/aptitude/skills/attention-concentration>
- Cognitive Skills Lab: attention practice:
<https://learn.novusstreamsolutions.com/cognitive-skills/attention-concentration>
- Emergency dispatch and 911 communications suite:
<https://learn.novusstreamsolutions.com/aptitude/tests/emergency-dispatch-and-911-communications>
- Error detection lesson: <https://learn.novusstreamsolutions.com/aptitude/skills/error-detection/lesson>
- Perceptual speed lesson: <https://learn.novusstreamsolutions.com/aptitude/skills/perceptual-speed/lesson>
- Attention and concentration lesson on Novus Learn:
<https://learn.novusstreamsolutions.com/aptitude/skills/attention-concentration/lesson>

Where this material comes from

- Character strings, invoice comparisons, batch counts and monitoring scenarios above are written for Novus Learn. All figures are original and no published test item is reproduced.
- Terminology and directions of effect follow widely published work: Mackworth (1948) on the vigilance decrement, Stroop (1935) on response conflict, Simons and Chabris (1999) on inattentional blindness, and the standard hit/miss/false-alarm framing from signal detection theory. Concepts only; no result figures are attributed to those studies.
- 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 not a clinical, diagnostic or attention-disorder screening measure.

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

Numerical reasoning

Arithmetic, fractions, percentages, ratios, rates, estimation, word problems, and number relationships.

Numerical reasoning is the ability to take a small set of supplied numbers (a price list, a staffing table, a fuel figure), and reach a defensible answer in around ninety seconds without a formula sheet. It is the most widely used quantitative section in graduate, management and public-sector screening, and it appears in banking, retail, armed forces and healthcare entry batteries alike. The arithmetic itself is deliberately ordinary: percentages, ratios, rates and averages. What is actually being measured is whether you pick the right

operation quickly, keep the units straight, and answer the question that was asked rather than the one you started computing. Practice here is device-local: no account, and nothing leaves this device unless you export it.

What you should be able to do after this lesson:

1. Move between fractions, decimals, percentages and ratios on sight, and pick whichever form makes a given item fastest (12.5 percent as one eighth, 33.3 percent as one third).
2. Apply percentage change in both directions, including recovering an original figure from a post-increase total, and combine successive changes multiplicatively rather than by adding them.
3. Split a total in a stated ratio by counting parts first, and work backwards from a stated difference between two shares.
4. Solve rate problems (speed, unit price, consumption, combined work rates) by carrying units through every intermediate line so the units of the answer prove the method.
5. Produce a one-significant-figure estimate before computing, and use it to eliminate the distractors built from the common wrong operations.
6. Read a multi-step stem and name, in one sentence, the single quantity the final question asks for before touching the numbers.

Worked examples and pitfalls

Reverse percentages: the item most candidates get backwards: A subscription price rose by 15 percent and now stands at 57.50 pounds. What was it before? The correct move is to divide, not subtract: the new price is 115 percent of the old, so the old price is $57.50 / 1.15 = 50.00$. Check it forwards: 15 percent of 50 is 7.50, and $50 + 7.50 = 57.50$. The tempting wrong answer takes 15 percent of the NEW figure, $0.15 \times 57.50 = 8.625$, and reports 48.88. That distractor is always on the option list because it is what most people do under time pressure, and it is wrong by 2.25 percent, small enough to look plausible. The same asymmetry drives the successive-change item: a price that rises 20 percent and then falls 20 percent does not return to where it started. Multiply the factors: $1.20 \times 0.80 = 0.96$, a net fall of 4 percent. Starting at 500 pounds you get 600 then 480, not 500. Percentages compose by multiplication; they never add.

Ratio splits: count the parts before you divide: A 4,830 pound budget is split between three teams in the ratio 3:5:6. Add the parts first: $3 + 5 + 6 = 14$. One part is $4,830 / 14 = 345$. The shares are $3 \times 345 = 1,035$, $5 \times 345 = 1,725$ and $6 \times 345 = 2,070$, and they sum back to 4,830, which is the check you should always run. Two classic failures. The first is dividing by 3 because there are three teams. That gives 1,610 each and ignores the ratio entirely. The second is treating 5 as a fraction and computing five sixths or five fourteenths of something other than the total. The more interesting version of this item gives you a difference instead of the total: 'the second team receives 690 pounds more than the first, what is the whole budget?' The difference between the shares is $5 - 3 = 2$ parts, so one part is $690 / 2 = 345$, and the budget is $14 \times 345 = 4,830$. Same number, reached from the other end, and the arithmetic is trivial once you have made the parts explicit.

Rates and units: 2 h 15 min is 2.25, not 2.15: A delivery van covers 174 km in 2 hours 15 minutes. Average speed is distance over time, and the time must be in hours: 15 minutes is $15/60 = 0.25$ h, so $174 / 2.25 = 77.3$ km/h. Type 2.15 into the calculator instead and you get 80.9 km/h: an error of roughly 4.7 percent that sits comfortably inside the plausible range and will match one of the options. Now extend it. The van consumes 8.6 litres per 100 km, so the trip needs $174 \times 8.6 / 100 = 14.964$ litres, and at 1.48 pounds per litre that is $14.964 \times 1.48 = 22.15$ pounds. Notice the units doing the work: (km) $\times$ (L / 100 km) leaves litres; (L) $\times$ (pounds / L) leaves pounds. If your intermediate line has km still attached at the end, you have divided when you should have multiplied. Write the unit next to every number and the method checks itself.

Unit pricing: normalise before you compare: Three pack sizes of the same fluid. Pack A: 750 ml for 3.60 pounds. Pack B: 2 litres for 9.40. Pack C: 1.5 litres for 7.20. Convert all three to price per litre. A: $3.60 / 0.75$

= 4.80 per litre. B: $9.40 / 2 = 4.70$. C: $7.20 / 1.5 = 4.80$. So B is cheapest by 10p a litre, and A and C are identical despite looking like different deals. The 'bigger pack is cheaper' heuristic happens to hold here but is not a rule, and test writers know it. Half of these items are built specifically so the largest pack loses. Now add the multibuy that these questions love: pack A is on three-for-two. Three packs give 2.25 litres for the price of two, 7.20 pounds, which is $7.20 / 2.25 = 3.20$ per litre and beats everything. The trap in the multibuy version is dividing by the number of packs paid for rather than the volume received.

Combined work rates: add rates, never times: Printer A completes a 3,000-page run in 50 minutes; printer B completes the same run in 75 minutes. Running together, how long? Convert to rates: A prints $3,000 / 50 = 60$ pages per minute, B prints $3,000 / 75 = 40$ pages per minute, and together they print 100 pages per minute, so the job takes $3,000 / 100 = 30$ minutes. Verify by counting output: in 30 minutes A produces 1,800 pages and B produces 1,200, which is 3,000 exactly. The two wrong answers you will see on the option list are 62.5 minutes (the average of 50 and 75) and 125 minutes (the sum). Both are impossible on inspection: two machines working together must finish faster than the faster machine alone, so any answer above 50 minutes is wrong before you compute anything. The general form is $1 / (1/50 + 1/75)$, and it is worth being able to write that line directly, but the sanity bound catches the error faster than the algebra does.

Estimate first, then compute: killing distractors in ten seconds: 'A department spent 847,300 pounds in 2024. In 2025 the budget fell by 12.5 percent. What was the 2025 spend?' Options: 741,387.50 / 953,212.50 / 105,912.50 / 762,570. Estimate before anything else: 12.5 percent is exactly one eighth, one eighth of roughly 850,000 is roughly 106,000, so the answer is roughly 744,000. That single line eliminates three options. 953,212.50 applied the change upwards. 105,912.50 is the size of the reduction, not the resulting spend: the 'answered the wrong question' distractor, and the most commonly selected wrong option on items of this shape. 762,570 is a 10 percent cut, planted for anyone who misread the rate. Exact working: $847,300 \times 7/8 = 5,931,100 / 8 = 741,387.50$. Doing it as a fraction avoids the decimal multiplication altogether. Learn the fraction equivalents cold (12.5 percent is $1/8$, 16.7 percent is $1/6$, 37.5 percent is $3/8$, 62.5 percent is $5/8$) because a fraction turns most percentage items into one division.

How to practise this skill

- Memorise the fraction-to-percentage table up to eighths and twelfths. Almost every 'nasty' percentage on these tests is a clean fraction in disguise, and the fraction route is two or three times faster than decimal multiplication.
- Write the unit beside every intermediate number, including the ones you keep in your head. Most numerical errors on timed sections are unit errors, not arithmetic errors, and units are the only cheap way to catch them.
- Rehearse reverse percentages as their own drill until dividing by 1.15 feels more natural than subtracting 15 percent. It is a single, identifiable technique that accounts for a disproportionate share of lost marks.
- Before selecting, re-read the last clause of the stem out loud. 'By how much did it fall' and 'what was it after the fall' have different answers, and both are always on the option list.
- Keep an error log with four columns (misread the question, wrong operation, unit slip, arithmetic slip), and tally which one you actually commit. Three sessions is usually enough to show that one column dominates, and that is the column to train.
- Work untimed until your method is stable, then add the clock in stages: generous, then realistic, then 20 percent tighter than the real test. Attempts stay on this device, so a slow first pass costs nothing.

Glossary

Percentage point: The arithmetic difference between two percentages. A rate moving from 4 percent to 5 percent rises by one percentage point but by 25 percent in relative terms; questions specify which they want and the two answers are both on the option list.

Reverse percentage: Recovering an original amount from a figure that already includes a percentage

change, by dividing by the multiplier (1.15 for a 15 percent rise, 0.88 for a 12 percent fall) rather than subtracting the percentage from the new figure.

Multiplier (decimal factor): The single number a quantity is multiplied by to apply a percentage change: 1.20 for plus 20 percent, 0.80 for minus 20 percent. Successive changes are handled by multiplying the factors, which is why plus 20 then minus 20 gives 0.96, not 1.

Ratio part: One unit of a ratio split. Dividing the total by the sum of the ratio terms gives the value of one part; every share and every difference between shares is then a whole number of parts.

Unit rate: A quantity expressed per one of something: price per litre, pages per minute, litres per 100 km. Comparisons are only valid once every option has been converted to the same unit rate.

Combined rate: The sum of two or more individual rates working simultaneously. Times cannot be added or averaged; only rates add, and the combined time is the total work divided by the combined rate.

Significant figure estimate: Rounding every input to one meaningful digit to get an approximate answer in a few seconds. Its purpose is not accuracy but elimination: it removes any option that is off by a factor or has the sign of the change wrong.

Distractor: A wrong option written deliberately to match a specific predictable error: subtracting instead of dividing, reporting the change instead of the result, averaging instead of combining rates. Recognising the pattern is often faster than recomputing.

Study this next

- Numerical reasoning skill hub: <https://learn.novusstreamsolutions.com/aptitude/skills/numerical-reasoning>
- Numerical reasoning practice in the Cognitive Skills Lab:
<https://learn.novusstreamsolutions.com/cognitive-skills/numerical-reasoning>
- Data interpretation lesson:
<https://learn.novusstreamsolutions.com/aptitude/skills/data-interpretation/lesson>
- Office numeracy suite: <https://learn.novusstreamsolutions.com/aptitude/tests/office-numeracy>
- Finance, accounting and insurance careers:
<https://learn.novusstreamsolutions.com/careers/families/finance-accounting-and-insurance>
- Numerical reasoning lesson on Novus Learn:
<https://learn.novusstreamsolutions.com/aptitude/skills/numerical-reasoning/lesson>

Where this material comes from

- All worked figures above were written and checked for Novus Learn. Every division, ratio split and rate calculation in this lesson was verified by recomputing it in the reverse direction.
- Conventions only (percentage point, unit rate, significant figures) cross-checked against standard public references such as the Wikipedia articles 'Percentage', 'Ratio' and 'Rate (mathematics)'. No item text is drawn from any published test.
- Novus Learn aptitude construct registry (catalog seed) for construct scope and suite mapping.
- Public educational framing only: not affiliated with any official exam board, publisher or employer, and no copyrighted test item is reproduced.

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

Verbal reasoning

Drawing conclusions from written information without relying on outside assumptions.

Verbal reasoning is the discipline of answering strictly from a passage: deciding whether a statement is True, False, or Cannot Say on the evidence given, and refusing to import anything you happen to know about the

topic. It is the workhorse construct in graduate and civil-service screening, banking and consulting sifts, and language-focused public-service batteries, where a passage about parcel volumes or grant conditions is followed by four statements to classify. The same discipline is what stops a contract review, a grant assessment or an incident summary from quietly acquiring facts nobody wrote down. Practice here is device-local: no account, and nothing leaves this device unless you export it.

What you should be able to do after this lesson:

1. Classify a statement as True, False, or Cannot Say, and state in one sentence which words in the passage force the classification.
2. Explain the difference that costs the most marks: 'Cannot Say' means undetermined by the passage, while 'False' means contradicted by it.
3. Spot quantifier drift between passage and statement (all, some, most, only, none) and show why 'all A are B' never licenses 'all B are A'.
4. Refuse a causal upgrade: recognise when a passage reports co-occurrence or sequence and a statement claims cause.
5. Separate what a passage asserts from what it reports somebody else asserting, and classify statements about each correctly.
6. Do the small arithmetic hidden inside verbal items - percentage changes, fractions of a stated total - without leaving the passage for outside data.

Worked examples and pitfalls

The three-way rule on one short passage: Passage: 'The Riverton depot handled 4,200 parcels in March, a 12 percent increase on February. A night shift was introduced in January; in March it processed 40 percent of the parcels the depot handled.' Now take four statements. (1) 'The depot handled 3,750 parcels in February.' February is March divided by 1.12: $4,200 / 1.12 = 3,750$ exactly, so this is True - the passage determines it even though the number is not printed. (2) 'The night shift caused the March increase.' Cannot Say. The passage puts the night shift and the increase in the same depot in the same period; it never claims one produced the other, and it never denies it either. Most candidates mark this False, reasoning correctly that correlation is not causation but then choosing the wrong label: False is reserved for statements the passage contradicts. (3) 'The night shift processed more than 1,700 parcels in March.' Forty percent of 4,200 is 1,680, which is not more than 1,700, so this is False - contradicted by arithmetic on the passage's own figures. (4) 'The depot handled more parcels in March than in January.' Cannot Say: January is mentioned only as the month the shift started, and no January volume is given.

Quantifier drift and illicit conversion: Passage: 'All accredited suppliers submit quarterly audits. Some suppliers in the northern region are accredited.' Statement A: 'Some suppliers in the northern region submit quarterly audits.' True, and it is worth seeing why the chain holds: at least one northern supplier is accredited, and every accredited supplier submits, so at least one northern supplier submits. Statement B: 'Every supplier that submits quarterly audits is accredited.' This reverses the first sentence. 'All accredited suppliers submit' leaves the door wide open for unaccredited suppliers to submit as well - perhaps voluntarily, perhaps under another rule - so the answer is Cannot Say. Turning 'all A are B' into 'all B are A' is called illicit conversion and it is the single most productive trap in this format, because the reversed sentence sounds like a paraphrase. Statement C: 'No unaccredited supplier submits quarterly audits' is the same reversal wearing a negative coat, and it is Cannot Say for the same reason. Statement D: 'All northern suppliers submit quarterly audits' is also Cannot Say: 'some are accredited' says nothing about the rest.

Asserted versus reported: Passage: 'The committee's report claims that the scheme cut waiting times by a third. Two of the four regional boards have disputed that figure.' Statement A: 'The scheme cut waiting times by a third.' Cannot Say. What the passage asserts is that a report claims this; the claim's truth is never settled, and the dispute does not settle it either. Statement B: 'At least two regional boards disagree with the report's waiting-time figure.' True, and note 'at least' - the passage says two disputed it, and two of four

disputing is consistent with 'at least two'. Statement C: 'All four regional boards accept the figure.' False, directly contradicted. Statement D: 'The report is wrong.' Cannot Say. This item type appears constantly in policy and diplomacy passages, where a paragraph stacks a claim, a source, and a reaction, and the reader has to keep three different truth-conditions apart. The reliable habit is to underline the reporting verb - claims, argues, estimates, alleges, projects - and remember that everything downstream of it belongs to the source, not to the passage.

Percentages are not symmetric: Passage: 'Membership fell from 8,000 to 6,000 over two years, then recovered by 25 percent in the third year.' Statement: 'Membership at the end of year three was higher than at the start.' The fall is 2,000 on a base of 8,000, which is 25 percent. The recovery is 25 percent of the new base: $6,000 \times 1.25 = 7,500$. That is below 8,000, so the statement is False. The trap is elegant, because a 25 percent fall followed by a 25 percent rise feels like it should cancel. It does not: to get back from 6,000 to 8,000 you need a 33.3 percent rise, since $2,000 / 6,000 = 0.333$. Whenever a verbal item states a change as a percentage, write down what the percentage is a percentage of before you touch it. A related version supplies the recovery as an absolute number instead - 'recovered by 2,000 members' - which does return the total to 8,000, and candidates who answered the first version from memory get the second one wrong.

Verbal analogies: name the relation as a sentence: Item: 'ANTISEPTIC is to INFECTION as ____ is to ____.' Options: (a) vaccine : immunity, (b) insulation : heat loss, (c) medicine : illness, (d) bandage : wound. Write the stem relation as a full sentence first: 'An antiseptic is applied in order to prevent an infection from occurring.' Now test each option against that exact sentence. Option (a) runs the other way - a vaccine is given to produce immunity, not to prevent it. Option (c) is the right family but the wrong specificity: medicine typically treats an illness that has already started. Option (d) is applied after the wound exists. Option (b) fits precisely: insulation is applied in order to prevent heat loss from occurring. The general method is the same on the simpler item types. 'BOOK is to LIBRARY as PAINTING is to ____' has canvas (the material), artist (the maker), and gallery (the place a collection is kept and shown) among its options; all three are genuine relations to 'painting', and only the third matches the stem sentence. Options in analogy items are chosen to be true statements about the words, just not the stated relation.

Two premises with 'some' prove nothing: Argument: 'Some depot managers are qualified engineers. Some qualified engineers hold a rigging certificate. Therefore some depot managers hold a rigging certificate.' This is invalid, and the way to prove it is a counter-model rather than an argument. Let the depot managers be Ana and Ben; let the qualified engineers be Ben and Cara; let the certificate holders be Cara alone. Premise one holds: Ben is a manager and an engineer. Premise two holds: Cara is an engineer with the certificate. The conclusion fails: neither Ana nor Ben holds the certificate. Since a single consistent world makes both premises true and the conclusion false, the argument cannot be valid, so in a True / False / Cannot Say framing the conclusion is Cannot Say. Two premises that both begin 'some' never combine into a conclusion, because 'some' gives you an overlap without telling you where the overlap sits. Contrast a valid pair: 'All accredited labs are inspected annually. No inspected facility is exempt from the fee.' Every accredited lab is inspected, and no inspected facility is exempt, so no accredited lab is exempt - True.

How to practise this skill

- Answer from the passage even when you know the subject. Candidates with a background in the topic score worse on verbal items than they expect, because their own knowledge quietly supplies the missing premise that turns a Cannot Say into a True.
- Run the two-worlds test on every candidate 'Cannot Say': can you imagine a world where all the passage's sentences hold and the statement is true, and another where they hold and it is false? If both, it is Cannot Say. If only the false world exists, it is False.
- Underline the quantifiers and hedges in the statement (all, some, only, most, may, must, likely) and find their counterparts in the passage. Roughly half of the wrong answers in this construct come from a single word swapped between the two.

- Budget about 25 to 30 seconds per statement rather than per passage, and read the passage once for structure before touching the statements. Re-reading the whole passage for each statement is what causes candidates to run out of time on the last set.
- Keep a wrong-answer log with one label per error: quantifier, causation, reported-versus-asserted, arithmetic, outside knowledge. A clear pattern almost always emerges within about forty items, and it is usually a single label.
- Work untimed until the three-way rule is automatic, then add the clock. Attempts are stored on this device only, so a slow first pass through a passage set costs you nothing but the time you spend on it.

Glossary

Entailment: A statement is entailed by a passage when it cannot be false while every sentence of the passage is true. Entailment is the only thing that earns a 'True' in this format; being plausible, likely, or well known does not.

Cannot Say: The verdict for a statement the passage neither entails nor contradicts. It is a claim about the passage, not about the world, which is why a statement you know to be true in real life can still be Cannot Say.

Quantifier: A word fixing how much of a group a claim covers: all, most, some, few, none, only. Swapping one quantifier for another changes the logical content completely while barely changing how the sentence reads.

Illicit conversion: The invalid move from 'all A are B' to 'all B are A', or from 'if P then Q' to 'if Q then P'. It preserves the words and destroys the logic, which is why converted sentences make such effective wrong answers.

Counter-model: A concrete, consistent scenario in which the premises hold and the conclusion fails. Producing one is the fastest possible proof that an argument is invalid, and it takes two or three named individuals.

Hedge: A qualifier such as may, could, is expected to, or is associated with. A hedged sentence in a passage cannot support an unhedged statement, and an unhedged passage sentence is not weakened by a hedged statement.

Reporting verb: A verb such as claims, argues, estimates, or alleges that attributes the following content to a source. Everything downstream of it is the source's assertion, and the passage takes no position on it.

Analogy relation: The specific link between the two stem words in an analogy item - tool and user, item and container, action and purpose. Naming it as a full sentence before reading the options removes almost all of the guesswork.

Study this next

- Verbal reasoning skill hub: <https://learn.novusstreamsolutions.com/aptitude/skills/verbal-reasoning>
- Practise the verbal reasoning bank:
<https://learn.novusstreamsolutions.com/cognitive-skills/verbal-reasoning>
- Critical thinking lesson: <https://learn.novusstreamsolutions.com/aptitude/skills/critical-thinking/lesson>
- Reading comprehension lesson:
<https://learn.novusstreamsolutions.com/aptitude/skills/reading-comprehension/lesson>
- General civil-service aptitude suite:
<https://learn.novusstreamsolutions.com/aptitude/tests/general-civil-service-aptitude>
- Verbal reasoning lesson on Novus Learn:
<https://learn.novusstreamsolutions.com/aptitude/skills/verbal-reasoning/lesson>

Where this material comes from

- Worked items written for Novus Learn. Every passage, statement set and figure above is original and invented for this lesson; no published or copyrighted test item is reproduced.
- Terminology follows standard, widely published usage in introductory logic and assessment writing - entailment, quantifier, illicit conversion, counter-model.
- Novus Learn aptitude construct registry (catalog seed) for the construct scope and the suite mapping shown in the related links.
- Public educational framing only - not affiliated with any official exam board, publisher or employer.

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

Error detection

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

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

What you should be able to do after this lesson:

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

Worked examples and pitfalls

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

A fixed scan order finds the one field that moved: Source record: Name Priyanka Ramaswamy / Account

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

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

Luhn: the number that fails, and the one case it misses: The Luhn check, used on payment and many membership numbers, doubles every second digit counting from the right, subtracts 9 from any doubled result above 9, sums everything, and requires a total ending in zero. Take 4539 1488 0343 6467. The doubled positions contribute 3, 3, 8, 0, 7, 2, 6 and 8, totalling 37; the undoubled positions contribute 7, 4, 3, 3, 8, 4, 9 and 5, totalling 43. The grand total is 80, which ends in zero, so the number passes. Now transpose the last two digits to 4539 1488 0343 6476. The doubled contributions become 5, 3, 8, 0, 7, 2, 6, $8 = 39$ and the undoubled become 6, 4, 3, 3, 8, 4, 9, $5 = 42$, giving 81. Not a multiple of ten, so the mistyped number is rejected at the point of entry. Luhn catches every single-digit substitution and almost every adjacent transposition: with exactly one blind spot. Swapping an adjacent 0 and 9 leaves the total unchanged, because doubling 0 gives 0 and doubling 9 gives 18 which reduces to 9, so both digits contribute the same amount whether doubled or not. A form that validates with Luhn will happily accept 90 where you typed 09. Knowing the blind spot is the point: a check digit narrows the space of undetected errors, it never closes it.

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

Reconcile the batch: consistent is not the same as correct: An invoice lists three lines: 12 at $14.25 = 171.00$; 7 at $33.60 = 235.20$; 3 at $128.00 = 384.00$. The stated subtotal is 709.20, VAT at 20 percent is stated as 141.84, and the total is stated as 851.04. Every downstream figure checks out against the one above it: $709.20 \times 0.20 = 141.84$ and $709.20 + 141.84 = 851.04$. Nothing in the tax or total column is inconsistent. But recompute the lines: $171.00 + 235.20 + 384.00 = 790.20$, not 709.20. The subtotal is a transposition of the correct figure, and because every later number was calculated from the wrong subtotal, the document is perfectly self-consistent and perfectly wrong. The true VAT should be 158.04 and the true total 948.24, a

shortfall of 97.20. This is the single most valuable habit in the construct: never verify a total against the number printed above it, always recompute it from the underlying items. Internal consistency proves that one person did the arithmetic carefully; it proves nothing about whether they started from the right number.

How to practise this skill

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

Glossary

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

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

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

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

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

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

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

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

Study this next

- Error detection skill hub: <https://learn.novusstreamsolutions.com/aptitude/skills/error-detection>
- Error detection practice in the Cognitive Skills Lab:
<https://learn.novusstreamsolutions.com/cognitive-skills/error-detection>

- Data checking and clerical accuracy lesson:
<https://learn.novusstreamsolutions.com/aptitude/skills/clerical-accuracy/lesson>
- Clerical checking suite: <https://learn.novusstreamsolutions.com/aptitude/tests/clerical-checking>
- Copyediting and proofreading suite:
<https://learn.novusstreamsolutions.com/aptitude/tests/copyediting-and-proofreading>
- Error detection lesson on Novus Learn:
<https://learn.novusstreamsolutions.com/aptitude/skills/error-detection/lesson>

Where this material comes from

- The ISBN-10 modulus 11 worked example (0-306-40615-2, weighted sum 132) and the Luhn worked example (4539 1488 0343 6467, sum 80) were computed by hand for this lesson and each was re-verified digit by digit, including the corrupted variants.
- Algorithm definitions cross-checked against the public specifications described in the Wikipedia articles 'International Standard Book Number' and 'Luhn algorithm', including the documented 09/90 transposition blind spot. Records, invoices and reference numbers above are invented.
- Novus Learn aptitude construct registry (catalog seed) for construct scope and suite mapping.
- Public educational framing only: not affiliated with any official exam board, publisher or employer, and no copyrighted test item is reproduced.

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

Recommended preparation

- Practice situational judgement:
<https://learn.novusstreamsolutions.com/aptitude/skills/situational-judgement/lesson>
- Strengthen critical thinking: <https://learn.novusstreamsolutions.com/aptitude/skills/critical-thinking/lesson>
- Practice attention and concentration:
<https://learn.novusstreamsolutions.com/aptitude/skills/attention-concentration/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/compliance-and-anti-money-laundering/study-outline>