

Forensic support roles: study guide

Police, justice, corrections, border, and security · suite apt-050-forensic-support-roles · generated 2026-09-15T15:27:10.105Z

Detail	Value
Title	Forensic support roles: study guide
Generated	2026-09-15T15:27:10.105Z
Fixture/version	apt-050-forensic-support-roles
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 Forensic support roles

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-050-forensic-support-roles&mode=guided>
- Start mini-test:
<https://learn.novusstreamsolutions.com/aptitude/practice/new?bank=apt-050-forensic-support-roles&mode=mini>
- Start full simulation:
<https://learn.novusstreamsolutions.com/aptitude/practice/new?bank=apt-050-forensic-support-roles&mode=full>

Skills covered, in full

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

Technical reasoning

Applied technical principles, diagrams, tools, systems, measurements, and troubleshooting.

Technical reasoning is what sits above any single trade: reading a system diagram for what it actually does, isolating a fault by measurement instead of by guesswork, and taking the governing number off a drawing or a nameplate without importing assumptions. It is assessed in HVAC, instrumentation, mechatronics, process-operator and engineering-technician selection, and it is the construct that best predicts whether someone can diagnose an unfamiliar machine. Employers care about it because part-swapping is expensive and half-splitting is not. Practice sessions here stay on your device unless you choose to export them.

What you should be able to do after this lesson:

1. Trace a process or signal diagram end to end and state, for a given symptom, which components could physically cause it and which could not.
2. Isolate a fault by half-splitting a chain of stages, and say how many measurements a chain of a given length should take.
3. Convert a dimension with asymmetric tolerance into an acceptance window and decide whether a measured part passes, can be reworked, or is scrap.
4. Match a measuring instrument to a required resolution, and distinguish an instrument's resolution from its accuracy.
5. Describe a closed control loop as sensor, controller, setpoint, actuator and feedback, then design the one measurement that separates a sensor fault from an actuator fault.
6. Read the qualifier attached to a rated number (duty cycle, working load limit, working pressure, nominal versus maximum), and apply it correctly.

Worked examples and pitfalls

Half-splitting beats swapping parts: A conveyor will not stop when its photo-eye is blocked. The chain is: sensor, field cable, junction box, PLC input card, PLC program, output card, interposing relay, contactor. Eight places the signal can die. Swapping parts one at a time means four or five attempts on average, since the culprit is equally likely to sit anywhere in the eight, and every attempt costs a part. Half-splitting starts in the middle instead: watch the PLC input LED while a colleague blocks the beam. If it toggles, the sensor, field cable, junction box and input card are all proven good in a single observation, and eight candidates become four. Force the output in the PLC and watch the contactor: if it pulls in, the output card, interposing relay and contactor are good as well, so the fault is in the program logic rather than the hardware. Three checks resolve eight stages, because each check halves the remaining suspects. The trap is starting at whichever end is easiest to reach, which resolves one stage per check instead of half of them. The second trap is the sentence 'I replaced the sensor and it still fails', that proves only that the sensor was not the fault, at the price of a part and an hour.

Tolerance: in spec, or scrap?: A drawing calls a shaft 25.00 mm with a tolerance of plus 0.05 and minus 0.10. That is an asymmetric tolerance, so the acceptance window runs from 24.90 mm to 25.05 mm and the nominal is not at its centre. A part measuring 24.92 mm is inside the window and passes, even though it is below nominal: the most common wrong rejection on this style of item, made by anyone who silently reads the tolerance as plus or minus 0.05. A part at 25.06 mm fails by 0.01 mm, but it fails oversize, so material can still be removed and it is rework rather than scrap. A part at 24.85 mm fails undersize and there is no recovering it. One more layer the better items include: if you took that 25.06 reading on a caliper with 0.02 mm resolution, the reading is at the very limit of what the instrument can resolve, and the honest next step is

to re-measure with a micrometer before anyone scraps or reworks anything.

Reading a system diagram: what can actually cause this?: A tank fill line is drawn as supply, isolation valve V1, strainer, pump P1, check valve, control valve CV1, tank. A high-level switch LSH-1 is wired to close CV1. The reported symptom is that the tank overfilled. Work the path between the measurement and the element that stops flow: a CV1 that has stuck open, an LSH-1 that never actuated, and a broken wire in the LSH-1 loop are all consistent with the symptom. A blocked strainer is not: restricting the inlet reduces flow, and no amount of restriction causes an overflow. Nor is the check valve, whose function is to prevent reverse flow, not forward flow. Candidates pick the strainer because it is the component they know fouls in service, which is a memory of maintenance history rather than a reading of the diagram. The discipline that earns the mark is directional: a component can only be responsible if it lies on the causal path AND its failure mode pushes the system in the direction of the symptom.

Instrument choice: resolution is not accuracy: A steel rule resolves to roughly 0.5 mm. A vernier caliper marked 0.02 mm resolves to 0.02 mm. A 0 to 25 mm micrometer resolves 0.01 mm on the thimble, or 0.001 mm if it carries a vernier. A dial indicator reads 0.01 mm of relative movement but tells you nothing about absolute size without a reference. Asked to verify a 25.00 mm shaft with a tolerance of plus or minus 0.02 mm, the tolerance band is 0.04 mm wide: two divisions on that caliper, which is not enough to judge anything reliably. The workshop convention is that the instrument should resolve to about a tenth of the tolerance band, here 0.004 mm, so even the micrometer is marginal and comparison against gauge blocks is the defensible answer. The trap the item is built around is a digital display: showing four decimal places is a statement about resolution, not accuracy. An uncalibrated digital caliper will report 25.0000 mm with total confidence and be 0.03 mm out.

Closed loop: which element failed?: A room is meant to hold 21 degrees Celsius. A thermostat containing the sensor and the controller drives a valve on a radiator. The symptom: the room reaches 28 degrees Celsius and the valve stays open. Three explanations survive first inspection. The sensor reads low so the controller still believes the room is cold, the valve is mechanically jammed open, or the controller output has failed in the on state. One measurement separates them. Put an independent thermometer beside the thermostat. If the thermostat displays 17 degrees while the thermometer reads 28, the sensor is lying and everything downstream is behaving correctly. If the thermostat displays 28 and is still calling for heat, the sensor is fine and the fault is in the controller or the valve, which you then split by checking whether the valve actuator is being energised. The tempting non-answer is 'the room is too hot, so lower the setpoint'. That treats the symptom, and if the sensor reads seven degrees low the loop will simply settle seven degrees high again at the new setpoint.

Nameplates: read the qualifier, not just the number: A welding machine is rated 200 A at 40 percent duty cycle over a ten-minute period. That means four minutes of arc time and six minutes of cooling in every ten, at the full 200 A. It does not mean 40 percent of 200 A, and it does not mean 40 percent of an hour. Both wrong readings feel entirely natural, which is why they make good distractors. Run the machine continuously at 200 A and the thermal cut-out will open. The same discipline transfers across the whole trade: a hoist's working load limit is not its breaking load, a hose's working pressure is not its burst pressure, a motor's service factor describes a short-term overload allowance and not a continuous rating, and a pump curve's flow figure is quoted at a stated head. Whenever an item hands you a number in a table or on a plate, underline the qualifier printed next to it before you calculate anything; the wrong options are usually built by dropping exactly one qualifier.

How to practise this skill

- Trace every diagram with a pencil from input to output and name each block as you pass it. Technical items punish skimming far harder than they punish slow arithmetic.
- Rehearse half-splitting on systems you already know (a home network that has dropped out, a car that will not crank), and count the checks. The habit transfers to the test intact.

- Underline the qualifier beside every number a question supplies: per hour, at 20 degrees Celsius, at 40 percent duty, nominal, maximum. That single mark-up defuses most distractors.
- Convert the whole question to one unit system in one pass before calculating, rather than converting each intermediate result and accumulating rounding.
- Train yourself to state what a measurement PROVES rather than what it reads. '12 V present at the coil' proves the supply and the whole upstream path; it says nothing about whether the coil itself is good.
- Keep a running list of the schematic symbols and component names you personally keep getting wrong, and drill only those. Five focused minutes beats another full untargeted set, and your attempt history stays local to this device so the list is yours alone.

Glossary

Half-split fault finding: Testing at the midpoint of a chain of stages so that each measurement eliminates half the remaining suspects. A chain of eight stages resolves in about three checks rather than four swaps.

Tolerance band: The distance between the upper and lower acceptance limits of a dimension. A tolerance written as plus 0.05 and minus 0.10 has a 0.15 mm band that is not centred on the nominal.

Resolution: The smallest change an instrument can display. A four-decimal readout has fine resolution and may still be inaccurate if the instrument is out of calibration.

Accuracy: How close a reading is to the true value, established by calibration against a traceable standard. Independent of resolution, and the property that decides whether a part passes.

Interlock: A condition wired or programmed to inhibit an action until it is satisfied, such as a guard-door switch that prevents a motor start while the door is open.

Closed-loop control: A control arrangement where a sensor measures the process, a controller compares that measurement to a setpoint, and an actuator drives the process until the difference closes.

Duty cycle: The proportion of a stated period during which equipment may operate at a stated output, for example 40 percent of a ten-minute period, after which it must cool.

Schematic versus pictorial diagram: A schematic shows function and connection logic with no regard to physical layout; a pictorial or exploded view shows physical arrangement and assembly order but hides the logic.

Root cause: The condition whose removal stops a failure recurring, as opposed to the symptom, which is only what became visible.

Study this next

- Technical reasoning skill hub: <https://learn.novusstreamsolutions.com/aptitude/skills/technical-reasoning>
- HVAC and refrigeration suite: <https://learn.novusstreamsolutions.com/aptitude/tests/hvac-and-refrigeration>
- Instrumentation and control suite:
<https://learn.novusstreamsolutions.com/aptitude/tests/instrumentation-and-control>
- Engineering and technical design careers:
<https://learn.novusstreamsolutions.com/careers/families/engineering-and-technical-design>
- Encyclopedia search: troubleshooting: <https://learn.novusstreamsolutions.com/search?q=Troubleshooting>
- Technical reasoning lesson on Novus Learn:
<https://learn.novusstreamsolutions.com/aptitude/skills/technical-reasoning/lesson>

Where this material comes from

- Diagnostic and metrology examples written for Novus Learn from general maintenance practice: binary-search fault isolation, asymmetric dimensional tolerance, and the ten-to-one instrument selection convention.

- Terminology checked against the public Wikipedia articles 'Troubleshooting', 'Engineering tolerance', 'Accuracy and precision' and 'Control loop'. Definitions only; every scenario above is original.
- Novus Learn aptitude construct registry (catalog seed) for the construct scope and related suite mapping.
- Public educational framing only: not affiliated with any official exam board, publisher or employer. Ratings and limits described here are illustrative and never override the equipment documentation in front of you.

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

Data checking and clerical accuracy

Matching, filing, coding, record checking, and identifying discrepancies.

Clerical accuracy is the ability to handle a large volume of records correctly and consistently: filing them in the right order, coding them against a key, spotting duplicates that do not look like duplicates, and holding that standard on the two-hundredth record as well as the second. Where error detection asks whether two things match, clerical accuracy asks whether you can apply a rule reliably at volume. It is assessed in administrative, records, clinical admin, school office, evidence-handling and insurance operations selection, and it is exactly the skill an employer is buying when they hire for a data-heavy back-office role. Practice stays on this device; there is no account and nothing is uploaded unless you export it.

What you should be able to do after this lesson:

1. Apply an alphabetical filing rule consistently, and state whether a given system files letter-by-letter or word-by-word. The two produce different, equally correct orders.
2. Sort alphanumeric references correctly under both string ordering and natural numeric ordering, and explain why record systems zero-pad.
3. Code records against a written key, handling every boundary condition (under, up to, and over, inclusive ranges) exactly as written rather than as intuited.
4. Identify duplicate records that differ only in formatting (case, whitespace, date order, punctuation inside a reference), and route genuinely ambiguous ones to exceptions instead of resolving them by guess.
5. Measure your own accuracy decay across a long batch and use the measurement to set a realistic attempt rate.
6. Convert a stated scoring rule into a target items-per-minute figure, and show why the optimum rate moves when the penalty multiplier changes.

Worked examples and pitfalls

Word-by-word or letter-by-letter: two correct answers, one rule: File these four names: Van Dyke, Vandenberg, Van Horn, Vance. Under word-by-word filing, each space-separated unit is compared in turn and 'nothing files before something', so the first unit 'Van' sorts ahead of both 'Vance' and 'Vandenberg'. The order is Van Dyke, Van Horn, Vance, Vandenberg. Under letter-by-letter filing, spaces are ignored entirely, so the keys are VANCE, VANDENBERG, VANDYKE, VANHORN, and the order is Vance, Vandenberg, Van Dyke, Van Horn. Both orders are correct filing; only one is correct for the system you are working in. Test items state the convention in the instructions, and the candidates who lose marks are almost always the ones applying whichever convention their previous employer used. The same fork appears with prefixes and punctuation: whether Mc and Mac interfile, whether St is treated as Saint, whether a hyphen counts as a space. Read the rule, restate it to yourself in one sentence, then apply it mechanically and do not let a name that 'obviously' belongs somewhere override it.

Alphanumeric codes: A-102 at the front or the back of the drawer: Sort the references A-7, A-12, A-70, A-102, B-3. Under natural numeric ordering, the way a person reads them, the answer is A-7, A-12, A-70, A-102, B-3. Under plain string ordering, the way most software sorts by default, each character is compared

in turn, so '1' comes before '7' and the answer is A-102, A-12, A-7, A-70, B-3, with A-102 first rather than last. Both are defensible; a filing item is testing which one the stated system uses. This is also why serious record systems zero-pad their references: rewrite the set as A-007, A-012, A-070, A-102 and the string sort and the numeric sort agree, permanently. If you are ever asked to design or clean a reference scheme, pad to a fixed width and the whole class of problem disappears. In a test, the giveaway is a set that deliberately mixes one-, two- and three-digit suffixes; that mixture exists only to separate candidates who apply the stated rule from candidates who apply the intuitive one.

Coding to a key: every mark is at a boundary: A claims routing key: under 500 pounds with no injury reported goes to CS1; under 500 with injury goes to CI2; 500 to 4,999 with no injury goes to CS3; 500 to 4,999 with injury goes to CI4; 5,000 and over goes to CE5 regardless of injury. Now code five records. R-118 at 499.99, no injury: CS1, since 499.99 is under 500. R-119 at exactly 500.00, no injury: CS3, because 'under 500' excludes 500 itself and the second band starts there. R-120 at 4,999.00 with injury: CI4, since the band is inclusive at its top. R-121 at exactly 5,000.00 with no injury: CE5, because the top band is 'and over' and its 'regardless of injury' clause overrides the injury split that governs the lower bands. R-122 at 86.40 with injury: CI2. Four of those five decisions turn on a boundary, and that is not an accident, coding items are written so that the interior cases are trivial and every discriminating mark sits on an edge. Before coding anything, underline the boundary words: under, up to, and over, between, inclusive. Then decide, once, what each one does to the endpoint, and apply that decision to every record in the batch.

Duplicates that are not textually identical: Three rows arrive in a merge. Row 1: SMITH, JANE | 07/04/1988 | AB123456C. Row 2: Smith, Jane | 1988-04-07 | AB 123456 C. Row 3: SMITH, JANE | 04/07/1988 | AB123456C. Rows 1 and 2 are the same person: normalise case, strip the spaces from the reference, and convert the ISO date and they match exactly. Row 3 is the genuinely hard one. If the file is in day/month order it is a different date of birth and possibly a different person; if that row came from a system using month/day order it is the same record again. You cannot tell from the row itself, so the correct action is to send it to the exception queue with the ambiguity noted, not to merge it and not to discard it. This is the judgement that separates competent records work from the appearance of it: the goal is not to make every row disappear, it is to make every decision defensible. In test form the item usually asks 'how many distinct individuals are represented', and the answer is often given as a range or accompanied by a 'cannot determine' option for exactly this reason.

Where accuracy actually decays on a long batch: Run a self-measurement rather than trusting a number from anyone: take 200 record pairs, split them into four blocks of 50, and log errors per block. Most people find block 1 slightly worse than block 2, a warm-up cost, and then a rise across blocks 3 and 4 as vigilance falls. The reason to measure it is that the arithmetic of small percentages is brutal at volume. Checking 200 pairs at 98 percent accuracy passes 4 bad records; at 99.5 percent it passes 1. Scale that to a realistic month of 20,000 records and the same two accuracy rates mean 400 defects against 100: a four-fold difference in downstream rework from a 1.5 point difference that would look like noise on a single test. Once you know where your own curve turns, the intervention is cheap: a deliberate twenty-second break at that point, or splitting the batch so the hardest records fall in your strongest block. Practice sessions on this site are recorded on your own device, so building a block-by-block picture across several sessions costs nothing but the logging.

Setting an attempt rate from the scoring rule: A 120-item checking test with a 10-minute limit, scored as correct minus incorrect. Suppose practice has told you that you hold 92 percent at ten items a minute and 96 percent at seven and a half. Fast: $10 \times 10 = 100$ attempted, 92 correct and 8 wrong, score 84. Careful: $10 \times 7.5 = 75$ attempted, 72 correct and 3 wrong, score 69. Fast wins by 15. Now change the rule to correct minus three times incorrect: fast scores $92 - 24 = 68$, careful scores $72 - 9 = 63$, and the 15-point gap has shrunk to 5. Now suppose your real accuracy at ten a minute is 80 percent rather than 92. A gap most people do not discover until they measure it. Fast now gets 80 right and 20 wrong: under simple correct-minus-incorrect that is 60, already behind the careful strategy's 69, and under the triple penalty it is $80 - 60 = 20$ against 63.

Same test, same person, opposite advice, and the two deciding variables are the penalty multiplier, which the instructions hand you, and your own accuracy-at-speed, which only measurement gives you. Do not pick a pace from temperament. Measure two rates in practice, write both accuracy figures down, and do this arithmetic before the test rather than during it.

How to practise this skill

- Before the first record of any batch, write the filing or coding rule at the top of the page in your own words. The single largest source of clerical error is applying a remembered rule from a previous system.
- Underline the boundary words in a coding key (under, up to, and over, inclusive), and resolve each endpoint once. Interior cases carry almost no marks; the edges carry nearly all of them.
- Normalise before you compare: mentally strip case, spaces and punctuation from references, and restate dates in one fixed order. Half of apparent duplicates are formatting, and half of apparent non-duplicates are too.
- When a record is genuinely ambiguous, mark it and move on. Time spent resolving one unresolvable row is taken from thirty rows you could have done correctly, and a guessed merge is worse than a flagged one.
- Measure your accuracy at two different speeds in practice and write both numbers down. You cannot choose an attempt rate rationally without them, and the fast rate is almost never as accurate as it feels.
- Practise on batches long enough to hit your own fatigue point, not on ten-item samples. The construct is specifically about sustained accuracy, and a ten-item drill measures the part of the curve that was never in doubt.

Glossary

Word-by-word filing: An alphabetical convention that compares space-separated units in turn, with a shorter first unit filing before a longer one. Under it, Van Horn files before Vance.

Letter-by-letter filing: An alphabetical convention that ignores spaces and punctuation entirely, comparing the run of letters. Under it, Vance files before Van Dyke.

Natural sort: Ordering that reads embedded digit runs as numbers, so A-7 precedes A-12. Contrasted with string ordering, which compares characters one at a time and puts A-102 first.

Zero padding: Writing references to a fixed width by adding leading zeros (A-007 rather than A-7) so that string ordering and numeric ordering produce the same result. The standard structural fix for alphanumeric filing errors.

Primary and secondary sort key: The field sorted on first and the field used to break ties within it. A batch sorted by site then by surname will look wrong if the two keys are applied in the opposite order.

Canonicalisation: Converting records to a single standard form (one case, no stray whitespace, one date order, punctuation stripped from references) before any comparison, so that formatting differences stop masquerading as data differences.

Exception queue: The destination for records that cannot be resolved from the information available. Routing an ambiguous record here is a correct outcome; guessing it into a merge is not.

Boundary condition: The endpoint of a coded band, where wording such as under, up to or and over decides which side a value falls. Coding tests concentrate their discriminating items here.

Study this next

- Data checking and clerical accuracy skill hub:
<https://learn.novusstreamsolutions.com/aptitude/skills/clerical-accuracy>
- Clerical accuracy practice in the Cognitive Skills Lab:

<https://learn.novusstreamsolutions.com/cognitive-skills/clerical-accuracy>

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

Where this material comes from

- All filing sets, reference codes, routing keys and records above were invented for this lesson. The two filing orders, the two sort orders and every score calculation in the attempt-rate example were worked through and checked by recomputation.
- Filing and sorting conventions cross-checked against standard public descriptions such as the Wikipedia articles 'Alphabetical order', 'Natural sort order' and 'Collation'. Terminology only; the examples are original.
- Novus Learn aptitude construct registry (catalog seed) for construct scope and suite mapping.
- Public educational framing only: not affiliated with any official exam board, publisher or employer, and no copyrighted test item is reproduced.

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

Safety judgment

Hazard recognition, safe sequencing, escalation, and risk controls.

Practise recognizing hazards and choosing a safe, proportionate response: protect people, control immediate exposure, use the correct reporting path, and verify that the control is effective.

What you should be able to do after this lesson:

1. Separate a hazard from the likelihood, severity, and exposure that shape its risk.
2. Choose an immediate control that stays within the person's training and authority.
3. Recognize when work should pause and when a supervisor, emergency process, or qualified specialist is needed.

Worked examples and pitfalls

Worked scenario: a damaged machine guard: A guard is loose before a scheduled run. Do not operate the machine or improvise a repair beyond your authorization. Keep people away, isolate or label the equipment only as procedure permits, report the defect to the responsible person, and wait for an approved inspection or repair. A deadline does not remove the hazard.

Risk-triage questions: Ask: what can cause harm, who is exposed now, how severe could the outcome be, what control is available, and who has authority to apply it? The safest answer is not always the most dramatic option; it is the option that controls the real exposure without creating a new hazard.

How to practise this skill

- Treat warnings, permits, isolation rules, and personal protective equipment as parts of a system, not interchangeable shortcuts.
- Do not ask an untrained person to investigate a hazard simply because they are nearby.
- In real work, site procedures and qualified safety direction override any general preparation heuristic.

Glossary

Hazard: A source or situation with the potential to cause harm.

Risk: A judgement about possible harm that considers likelihood, severity, and exposure.

Control: A measure that removes a hazard or reduces exposure to it under an approved process.

Study this next

- Occupational safety and health:
<https://learn.novusstreamsolutions.com/search?q=Occupational%20safety%20and%20health>
- Hazard: <https://learn.novusstreamsolutions.com/search?q=Hazard>
- Risk assessment: <https://learn.novusstreamsolutions.com/search?q=Risk%20assessment>
- Safety judgment lesson on Novus Learn:
<https://learn.novusstreamsolutions.com/aptitude/skills/safety-judgement/lesson>

Where this material comes from

- Novus Learn original safety-critical, hazard-recognition, and operational suite scenarios.
- Novus educational framework: identify the hazard, protect people, control exposure, report, and verify.

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

- Learn safety judgement: <https://learn.novusstreamsolutions.com/aptitude/skills/safety-judgement/lesson>
- Strengthen clerical accuracy:
<https://learn.novusstreamsolutions.com/aptitude/skills/clerical-accuracy/lesson>
- Build 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/forensic-support-roles/study-outline>