



Level 1 — The Common Core

Twelve hours, every cadre, taught together. Facilitator edition.

Dr Neal Aggarwal

Written in a personal capacity

AUGUST 2026

*Annex C to Another Arrow in the Quiver. Facilitator edition, version 0.1 — **not yet piloted.***

I write here in a personal capacity.

How to use this document

This is the teaching material for the twelve-hour Foundation course, written so that a certified instructor can run it from the page. Each unit gives you learning outcomes mapped to the competency codes in Annex B, a timed session plan, the facilitator notes that matter, the learner activities in full, and the assessment specification.

What this document is not. It has not been piloted. Every module in the Institute’s catalogue is required to run against a pilot cohort with think-aloud observation before it enters delivery, with pilot data going to the Office of Quality rather than to the authors. This version has not been through that. Treat the timings in particular as estimates that the pilot will correct — in my experience the discussion-heavy units always overrun and the didactic ones always finish early.

Design decisions you should understand before you teach it

Mixed cadre, deliberately. The consultant surgeon and the ward nurse sit in the same room for all twelve hours. This is not an efficiency measure. The failure modes this course exists to prevent live in the space *between* cadres — the nurse who notices that an AI-supported decision looks wrong and does not say so, the consultant who does not expect to be challenged. A nurse who has been taught to escalate needs to have practised doing it in front of a consultant who has been taught to expect it. If you run cadre-segregated cohorts you will have a smoother day and a worse course.

Groups of 16 to 24. Below 16 the adversarial exercises lose the range of clinical experience they depend on. Above 24 the facilitator cannot observe individual performance in Unit 1.5, which is the unit that matters most.

Two facilitators. One leading, one observing and noting. The observation is not optional — Unit 1.5 generates the individual performance data that the Office of Quality uses to calibrate the assessment.

Order is fixed. Unit 1.5 must follow 1.1–1.4 and must precede 1.6–1.7. Learners cannot practise discernment before they have a mental model of why these systems fail, and should not touch a real workflow before they know the law.

Prerequisites and materials

Requirement	Detail
Learner prerequisite	Current registration with KMPDC, NCK or the relevant council. No prior AI experience assumed or expected
Facilitators	Two, both Level 4 certified. At least one from a clinical cadre represented in the room
Room	Flat floor, movable seating for small groups. Not a lecture theatre
Devices	One laptop or tablet per two learners, on the Institute sandbox. Personal devices are not used
Sandbox	Seeded-error configuration active from Unit 1.5. Interaction logging on, with learner consent taken at registration
Printed	Case packs 1–4, non-delegable list worksheet, personal AI policy template
Time	12 contact hours. Two days of six, or three of four. Not one long day — Unit 1.5 needs learners who are not tired

Assessment at a glance

Instrument	Covers	Weight
Invigilated MCQ, 40 items from the blueprinted bank	D1.F1–F4, D2.F1–F4, D3.F1–F5, D4.F1–F4	Pass/fail, standard set by modified Angoff
Observed simulated encounter, one station	D2.F1, D3.F1, D3.F3	Pass/fail, conjunctive
Personal AI policy, written	D1.F2, D4.F1, D4.F4	Satisfactory / not yet satisfactory

Two conjunctive requirements. A learner who enters identifiable patient data into an uncontrolled system in the simulated encounter (D2.F1), or who fails to detect the seeded error at threshold (D3.F3), fails the course regardless of aggregate score. Tell them this at the start of Unit 1.1. It changes how they attend.

Unit 1.1 — What these systems are and are not

2.0 hours · plenary and paired discussion

Learning outcomes

By the end of this unit the learner can:

1. Describe in plain language how a large language model produces text, without technical detail beyond what is needed for clinical judgement. (*underpins D3.F4*)

2. Explain why fluency and accuracy are independent properties of an output. (*D3.F4*)
3. Distinguish a language model from a diagnostic classifier and from a rule-based decision support alert, and state why the distinction changes how much they should trust each. (*D1.F1*)
4. Give three examples of failure modes these systems exhibit in Kenyan clinical practice. (*D3.P2, introduced*)

Session plan

Time	Activity
0:00	Welcome. The two conjunctive requirements stated explicitly
0:10	Opening exercise: in pairs, “what have you already used one of these for, in the last month?” Collect on the board without comment or judgement
0:30	Facilitator input: how next-token prediction produces fluent text
0:55	Demonstration: the fabricated citation (see below)
1:15	Break
1:25	Three kinds of system: language model, classifier, rule-based alert
1:50	Paired task: sort eight described tools into the three categories
2:00	Close

Facilitator notes

The opening exercise sets the tone for the whole course and you can ruin it in thirty seconds. People will admit to things — pasting a patient summary into a consumer chatbot, asking for a dose at 2 a.m. Write it up without flinching. If anyone is made to feel foolish in the first half hour you will get nothing honest for the remaining eleven and a half hours, and the course depends on honesty about current practice.

On the technical explanation. Give them enough to support judgement and no more. The useful mental model is: the system produces the most plausible *continuation*, and plausibility is a property of the text, not of the world. That single idea does most of the work in this course. Resist the temptation to teach architecture; you are not making machine-learning engineers, and a room of clinicians will disengage.

The demonstration. Ask the sandbox for the evidence on a specific clinical question in a Kenyan context, and request references. It will typically produce references in correct citation format, some of which will not exist. Look one up live. Do not pre-select an example that fails — run it live and take what you get. If it happens to produce entirely genuine references, that is a *better* teaching moment than a rigged one: point out that you could not know which case you were in without checking, which is exactly the problem.

Anticipated difficulty. Someone will say “so it is just autocomplete, why would I use it at all?” That is the under-trust failure mode arriving early. Do not crush it — note it, and say you will return to it in Unit 1.5. Then make sure you do.

Assessment items (specimen)

Item 1.1.a — A registrar asks a language model for the dose of a drug in renal impairment. The system gives a confident, precisely-worded answer with a plausible reference. Which of the following is the best account of what has happened?

A. The system has retrieved the dose from a drug database B. The system has produced the most plausible continuation of the text, which may or may not correspond to the correct dose ✓ C. The system has calculated the dose from the patient’s parameters D. The confidence of the wording indicates the answer has been verified

Rationale: tests outcome 1 and 2 together. Distractor D is the important one — it is the automation-bias trap in miniature, and item analysis should watch whether strong candidates select it.

Item 1.1.b — Which of the following is a rule-based alert rather than an AI system as defined in the competency standard?

A. A system suggesting a differential diagnosis from a free-text history B. A system flagging every prescription of gentamicin where the recorded creatinine exceeds a set threshold ✓ C. A system estimating the probability of sepsis from vital signs using a trained model D. A system drafting a discharge summary from structured notes

Unit 1.2 — The Clinical 4Ds

2.0 hours · plenary, small group

Learning outcomes

1. Name the four competencies and state the question each answers. (*D1–D4*)
2. Distinguish automation, augmentation and agency, and identify which they are operating in for a given task. (*D1.F1*)
3. Describe the two loops — Delegation/Diligence outer, Description/Discernment inner — and locate a clinical encounter within them.

Session plan

Time	Activity
0:00	The four domains, with the clinical question each answers
0:25	The three modalities, with a clinical instance of each
0:50	Small group task: modality sort (Case Pack 1)
1:15	Break
1:25	The two loops, and why clinicians already recognise the shape
1:45	Plenary: where does each modality fail?
2:00	Close

Facilitator notes

The modality distinction is the highest-yield twenty minutes in the entire course. Almost every serious failure I can construct involves someone operating in *agency* mode while believing they are in *augmentation* mode — a triage assistant configured once and left running on a queue, with nobody in the room when it goes wrong. Spend the time.

On the two loops. Delegation and Diligence form the outer loop: what you decide to hand over, and what you answer for afterwards. Description and Discernment form the inner loop, turning many times within a single encounter. Clinicians recognise this immediately when you name it as consent-and-audit on the outside, history-and-examination on the inside. Use that mapping; it lands every time.

Case Pack 1 — modality sort

For each, the group states the modality and the dominant risk.

1. A clinical officer asks a system to draft a referral letter from her own notes. (*Automation — unreviewed output entering the record*)

2. A physician works through a difficult differential with a system, revising as he goes. (*Augmentation — anchoring; his own reasoning quietly revised*)
 3. A hospital configures a triage assistant to pre-sort the outpatient queue each morning. (*Agency — harm at scale, no clinician present when it occurs*)
 4. A midwife asks a system to explain pre-eclampsia to a patient in Kiswahili. (*Automation, with a translation-quality risk she may not be able to check*)
 5. A surgeon asks for the evidence on a technique, then argues with the answer. (*Augmentation — the healthiest pattern in the pack; say so*)
-

Unit 1.3 — Delegation and the non-delegable list

1.5 hours · individual work, then structured challenge

Learning outcomes

1. Construct a personal non-delegable list and defend each entry. (*D1.F2*)
2. Explain why reversibility and severity of error, not convenience, determine delegability. (*D1.F3*)
3. Decline AI use where they lack the knowledge to evaluate the output. (*D1.F4*)

Session plan

Time	Activity
0:00	The concept. My own list offered — as a starting point for argument, not as doctrine
0:15	Individual: draft your own list on the worksheet
0:35	Pairs, cross-cadre: exchange lists and challenge every entry
1:00	Plenary: where did lists differ by cadre, and why?
1:20	The reversibility/severity test
1:30	Close

Facilitator notes

Pair across cadres in the challenge phase. A nurse challenging a surgeon's list, and vice versa, is the point of the exercise. The differences are instructive: nursing lists tend to be longer and more absolute; medical lists tend to be shorter with more conditionals. Neither is wrong, and the discussion of *why* they differ teaches more than the lists do.

Do not present my list as the answer. It is: consent, breaking significant news, the decision to operate, the final diagnosis committed to the record, the prescription, the signature. If a learner argues one of those off their list with a good reason, that is a better outcome than compliance. What you are assessing is whether they can *defend* a boundary, not whether they picked yours.

Watch for the convenience trap. Someone will justify delegating a task because they are busy. Name it: workload is a real pressure and a bad reason. The test is what happens if the output is wrong and nobody notices.

Worksheet — the non-delegable list

For each act, state: (a) would I hand this to a system? (b) if not, why not — what is the error I am protecting against? (c) what would have to change for my answer to change?

Obtaining consent · Breaking significant news · Deciding to operate · Committing the final diagnosis to the record · Writing the prescription · Signing the record · Drafting a discharge summary · Generating a differential · Summarising the literature · Translating patient information · Coding for billing · Triaging a queue

Unit 1.4 — Description as clinical handover

1.5 hours · demonstration and practice

Learning outcomes

1. Structure a clinical query using a recognised handover structure. (D2.F3)
2. Specify facility level, available investigations, formulary and local epidemiology. (D2.F2)
3. Request reasoning, differentials and the case against. (D2.F4)
4. State and apply the absolute rule on identifiable data. (D2.F1 — conjunctive)

Session plan

Time	Activity
0:00	The absolute rule, stated first and unambiguously
0:10	A good prompt is a good handover: SBAR mapped onto query structure
0:30	Demonstration: the same question asked badly, then well
0:50	Paired practice on the sandbox, Case Pack 2
1:20	Plenary: what changed in the answers?
1:30	Close

Facilitator notes

Open with the rule, not with technique. Identifiable patient data does not go into a system you or your employer do not control. Say it, write it up, and say that it is one of the two things that will fail them outright. Then teach the craft.

The comparison demonstration does the teaching. Ask a bare question — “management of severe pre-eclampsia” — and take the answer. Then ask it properly: a Level 4 facility, magnesium sulphate available, no ICU, no on-site obstetrician overnight, this is a referral decision. The difference is large and self-evident. You will not need to argue for structure after they have seen it.

The most common error in practice is not poor phrasing. It is omitting the context a Kenyan clinician takes so completely for granted that it does not occur to them to state it — which investigations actually exist at this facility, what is on the formulary this month, what the local prevalence of the thing they are worried about actually is. The system defaults to a well-resourced setting unless told otherwise.

Language. A meaningful proportion of encounters are not conducted in English. Where a history was taken in Kiswahili or Dholuo and rendered into English for the query, that translation is lossy and the system cannot signal that it is reasoning about a lossy trans-

lation. Ask the room how they would flag it. There is no clean answer; the point is that they notice.

Case Pack 2 — paired practice

Each pair takes one case, drafts a query, runs it, then improves it.

2.1 A 24-year-old primigravida at 34 weeks, BP 168/112, 3+ proteinuria, headache. Level 4 sub-county hospital. Magnesium sulphate in stock, no ICU, nearest obstetrician 90 minutes away.

2.2 A 6-year-old, three days of fever, no rash, RDT negative for malaria, mildly reduced consciousness. Level 3 health centre. No laboratory after 4 p.m.

2.3 A 58-year-old man on gentamicin day 4, creatinine risen from 96 to 180 $\mu\text{mol/L}$, still febrile. Level 5 county referral hospital, no therapeutic drug monitoring.

2.4 A 31-year-old woman, known HIV on dolutegravir-based ART, now pregnant, asking about safety. Level 4 facility.

Note for facilitators: all four are constructed and contain no real patient data. Learners should be told this explicitly, and told that constructing such a case is itself a skill — it is how they will ask questions safely in practice.

Unit 1.5 — Discernment I: the independent-impression rule

2.0 hours · the unit the course exists for

This is the unit that justifies the other ten hours. If you run out of time elsewhere, cut elsewhere.

Learning outcomes

1. Form and record an independent clinical impression **before** consulting the system. (D3.F1)
2. Read the reasoning before the conclusion. (D3.F2)
3. Detect plausible, fluent and wrong output in a clinical vignette at threshold. (D3.F3 — *conjunctive*)
4. Verify citations, doses, thresholds and guideline references before acting. (D3.F5)
5. Articulate calibrated trust, and recognise both directions of miscalibration. (D3.P3, *introduced*)

Session plan

Time	Activity
0:00	The evidence: the 2025 trial, stated plainly
0:15	Calibrated trust — the inverted U. Under-trust is a real failure, not a safe default
0:30	The independent-impression rule. Drill it
0:40	Seeded-error exercise, round 1 (Case Pack 3, three cases)
1:10	Break
1:20	Debrief round 1: who caught what, and <i>when</i> did they know?
1:40	Seeded-error exercise, round 2 (two cases, rate not disclosed)
1:55	Close: what you will do differently on Monday
2:00	End

Facilitator notes

Open with the evidence and do not soften it. In 2025 a randomised trial gave clinical cases to physicians who had *already completed twenty hours of AI-literacy training*, and exposed half of them to deliberately erroneous model output. They deferred to it. Consultation was voluntary; they retained full autonomy to reject the suggestion. They accepted it.

Say directly: *the training you are receiving today does not, on its own, protect you. That is why the rest of this unit is a drill and not a lecture.* Learners take the exercise far more seriously once they understand the course is not claiming to inoculate them.

The independent-impression rule is the single highest-value behaviour in the course. *Form and record your own clinical impression before you look at the output.* Not after. It makes the clinician's own reasoning a fixed point rather than something that gets quietly revised in the light of a confident answer. In the exercise it is mechanical: they write their impression on the card, face down, before the sandbox is unlocked.

On the seeded-error rate. Round 1 runs at one in three and you tell them the rate. Round 2 runs at an undisclosed rate — it may be zero, one or two of the cases. Do not tell them afterwards how many there were until the debrief, and note who searched for errors that were not there. Over-detection is its own miscalibration and belongs in the discussion.

The errors are not cartoonish. They are the errors these systems actually make: a plausible dose wrong for renal impairment; a differential omitting the tropical diagnosis; a confident citation to a paper that does not exist; a guideline recommendation correct for a European population and wrong here; a drug name right internationally that refers to something else locally.

The debrief question that does the teaching is not “who caught it?” but “**at what point did you know?**” Some will have felt something was wrong and continued anyway. That is the automation-bias moment made visible, and it is worth more than the score.

Watch for the learner who catches nothing and is unbothered. Flag them to the observing facilitator. They are not failing the exercise; they are demonstrating the disposition the course exists to change, and they need the Practitioner level more than anyone else in the room.

Case Pack 3 — seeded-error exercise, round 1

For each: learners write their own impression, face down. Then the sandbox output is released. Then they mark: *do I accept, modify, or reject — and why?*

3.1 — Seeded: dose error in renal impairment. A 62-year-old man, community-acquired pneumonia, creatinine 210 µmol/L. Output recommends a standard antibiotic regimen at full dose with no renal adjustment, in confident and well-organised prose. *What we are testing:* D3.F5. Roughly two-thirds of learners catch this. Those who do not are usually the ones who did not write an impression first.

3.2 — Seeded: omitted differential. A 29-year-old woman, two weeks of fever, night sweats, weight loss, non-productive cough, HIV status unknown. Output gives a competent differential centred on atypical pneumonia and does not mention tuberculosis. *What we are testing:* D3.P2(a) — under-weighting of locally prevalent conditions. This one produces the most useful discussion in the whole course. Almost every Kenyan clinician catches it instantly, which is exactly the point: **their local knowledge is the safeguard, and the system does not have it.**

3.3 — Not seeded. The output is sound. A 45-year-old with newly diagnosed type 2 diabetes, no complications. Output gives appropriate, well-structured initial management. *What we are testing:* over-detection. Some learners will invent a fault. Use this in the debrief to introduce under-trust as a genuine failure mode with a real cost — the clinician who refuses the screening tool and misses the retinopathy has harmed a patient just as surely as the one who accepted a wrong answer.

Case Pack 4 — round 2, rate undisclosed

4.1 A 3-year-old with severe acute malnutrition and diarrhoea. *Seeded: fluid resuscitation guidance appropriate for a well-nourished child, which is dangerous in SAM.*

4.2 A 70-year-old on warfarin with a new prescription. *Not seeded; the interaction is correctly flagged.*

Unit 1.6 — Law, consent and documentation

1.5 hours · plenary with worked scenarios

Learning outcomes

1. State the requirements of the Data Protection Act 2019 and the Digital Health Act 2023 as they bear on entering, storing and transferring health data. (D4.F2)
2. Document in the clinical record where AI materially informed a decision, and how. (D4.F1)
3. Disclose AI involvement to a patient where a reasonable patient would wish to know. (D4.F3)
4. State that accountability remains wholly with the registrant. (D4.F4)

Session plan

Time	Activity
0:00	The legal frame: Data Protection Act 2019, Digital Health Act 2023
0:25	What may and may not be entered — worked through, not asserted
0:45	Documentation: what a good record entry looks like
1:05	Disclosure to patients: role-play in pairs
1:25	Accountability: the signature is yours
1:30	Close

Facilitator notes

Teach the principle, not the section numbers. Clinicians will not retain statutory citations and do not need to. What they need is the operating rule — identifiable data stays inside systems you or your employer control — and the knowledge that a statutory frame exists and that their employer’s data protection officer is the person to ask.

On documentation, give them a form of words. Most learners have never seen one. Something like: *“Differential generated with AI decision support; reviewed and modified by me; final impression and plan as above are my own.”* The point is that it is specific about what the system did and unambiguous about who decided.

The disclosure role-play is harder than it looks and learners consistently underestimate it. One plays a patient who asks, directly, “did a computer decide this?” The honest answer is neither “no” nor “yes” — it is that a system was used to help think it through, the clinician reviewed it, and the decision is theirs. Watch for learners who become defensive; that is the tell that they have not internalised D4.F4.

Anticipated question you must not fudge: *“if I follow the AI and it is wrong, am I covered?”* No. The signature is yours, the liability is yours, and no tool has ever taken responsibility for anything. Say it plainly.

Unit 1.7 — First supervised use

1.5 hours · sandbox, observed

Learning outcomes

1. Complete a full clinical query cycle applying all four domains. (*D1-D4*)
2. Produce a written personal AI policy. (*D1.F2, D4.F1, D4.F4*)

Session plan

Time	Activity
0:00	Brief: bring a real question from your own practice, de-identified
0:10	Supervised sandbox work , facilitators circulating and observing
0:50	Write your personal AI policy (template provided)
1:10	Pairs: read each other's policy and challenge one clause
1:25	Close: what happens next, and the recertification expectation
1:30	End

Facilitator notes

Circulate and watch hands, not screens. What you are looking for is whether the independent-impression rule survived four hours after Unit 1.5. My working hypothesis — which I would want tested and would not be surprised to see refuted — is that it is the first discipline to decay. Note who writes their impression first without being reminded; that is the observation the Office of Quality most wants.

The personal AI policy is the artefact the learner leaves with. It should fit on one side of a page. It is not assessed for elegance; it is assessed for whether the learner can state a boundary and a documentation practice they will actually keep.

Personal AI policy — template

1. What I will use AI for in my practice **2. What I will never delegate** (*your non-delegable list*) **3. What I will never enter into a system I do not control**
4. How I will check output before acting on it **5. How I will record AI-assisted decisions** **6. What I will do if I think AI contributed to a harm or near miss** **7. When I will review this policy**

Assessment specification

Knowledge assessment

40 items, invigilated, closed-book. Blueprinted across the domains:

Domain	Items	Proportion
D1 Delegation	8	20%
D2 Description	8	20%
D3 Discernment	16	40%
D4 Diligence	8	20%

The 40% weighting on Discernment is deliberate and should not be traded away for a more even distribution. It reflects where the evidence says the risk sits.

Standard setting by modified Angoff, panel of not fewer than eight practising clinicians spanning the cadres in the cohort. Never an arbitrary 50%; never a pass rate agreed in advance.

Item analysis after every sitting. Difficulty index, discrimination index, distractor analysis. Items with negative discrimination are withdrawn immediately and affected candidates' scores recalculated. Reliability reported and published.

Simulated encounter

One station, ten minutes, standardised patient, sandbox available. Seeded error present. Scored on a structured rubric across four domains:

Domain scored	Weight	Conjunctive?
Appropriate delegation	20%	No
Quality of description	20%	No — except identifiable-data breach, which fails outright
Detection and correction of error	40%	Yes
Documentation and disclosure	20%	No

Examiner calibration before every diet, with inter-rater reliability monitored across stations and examiners. An examiner whose scoring drifts is retrained before examining again.

Personal AI policy

Marked satisfactory / not yet satisfactory against three criteria: a defensible non-delegable list; a specific and workable documentation practice; and an unambiguous statement of personal accountability. “Not yet satisfactory” triggers a rewrite, not a fail.

Award and progression

Learners meeting all three requirements receive the **Certificate in Clinical AI Foundations**, with CPD points at the rate agreed with their Council.

The certificate is a prerequisite for Level 2 Practitioner in any track. It is **not** a licence to use AI in patient care unsupervised; Level 2, with its countersigned workplace logbook of twenty supervised encounters, is where that is established.

Recertification at not more than two years, assessed rather than attested.

What we do not yet know about this course

Written into the material deliberately, because a facilitator teaching it should know what is uncertain.

Whether twelve hours is the right length. It is a judgement. The pilot may show that Unit 1.5 needs to be twice as long and Unit 1.1 half.

Whether the seeded-error rate of one in three is right. Too low and learners do not encounter enough errors to build the reflex; too high and they learn to distrust everything, which is the under-trust failure. One in three is a starting point, not a finding.

Whether the independent-impression rule survives contact with real work. This is the most important open question in the whole programme. The Institute measures it at three and twelve months, and will publish what it finds.

Whether mixed-cadre teaching works as well as I believe it does. It is a strong design claim on thin evidence. If the pilot shows that nurses speak less in mixed rooms rather than more, the design is wrong and must change.

Prepared by Dr Neal Aggarwal, August 2026, in a personal capacity. Version 0.1, not yet piloted. Drafted with AI assistance; the pedagogical judgements and any errors are mine.

The four domains adapt the AI Fluency Framework of Prof. Rick Dakan and Prof. Joseph Feller, elaborated into an open course series with Anthropic PBC. Course materials are CC BY-NC-SA 4.0; this derivative carries that licence forward. See Annex B §13.