



C O M P U T E R F O R E N S I C S L A B

London digital forensics & e-discovery · Established 2007

Using Generative AI Responsibly in Electronic Disclosure

Prompted Review, Human Verification and Defensibility Under PD 57AD

Generative AI has moved from experiment to working tool in UK disclosure: prompted models now assist relevance assessment, privilege flagging, summarisation and redaction support in court-ordered exercises under PD 57AD, and best-practice guidance has followed: the ILTA Generative AI Best Practice Guide for outgoing disclosure (September 2025), covered by the Law Society's practical guidance in November 2025, alongside judicial AI guidance and the accountability line drawn by the Divisional Court in *Ayinde*. The technology's promise is real: reading at scale, consistent application of instructions, cost profiles that reshape proportionality. So are its failure modes: confident errors, prompt sensitivity, outputs that vary between runs. This guide sets out responsible use for UK lawyers: where GenAI fits in the disclosure workflow, the prompt-testing and verification disciplines that make it defensible, the accountability and transparency framework, and the protocol drafting that keeps the machine's reading inside the lawyer's duty.

🕒 Estimated reading time: 25 min read

§ ABOUT THE AUTHOR

PREPARED BY COMPUTER FORENSICS LAB E-DISCOVERY TEAM

Computer Forensics Lab is a London-based digital forensics and electronic disclosure practice established in 2007. Its eDiscovery arm, **e-discovery.uk**, deploys generative AI in disclosure workflows under the verification and audit disciplines this guide describes: prompt engineering and testing, sampled validation against human benchmarks, and the documentation that makes AI-assisted exercises defensible under PD 57AD: with the same honest scoping it applies to every technology: including when not to use it.

ESTABLISHED 2007 · LONDON

ISO 17025-ALIGNED PROCEDURES

GENAI-ASSISTED REVIEW WORKFLOWS

VALIDATION & AUDIT DESIGN

CPR PART 35 EXPERT REPORTS

FULL CHAIN-OF-CUSTODY DOCUMENTATION

§ CONTENTS

In this guide

01	Executive summary
02	The problem in plain English: a reader that never tires and sometimes invents
03	The guidance landscape: ILTA, the Law Society, the judiciary and Ayinde
04	Where GenAI fits: relevance, privilege, summarisation and redaction support
05	Prompt engineering as methodology: drafting, testing and versioning
06	Verification: sampling AI outputs against human judgement
07	Transparency, data protection and the protocol
08	Source architecture: where else the evidence lives
09	Worked examples
10	Common mistakes and technical limitations
11	Questions to ask · Suggested wording
12	Checklist and red flags · When to involve a digital forensic expert
13	Frequently asked questions
14	Glossary · References · Disclaimer · How a specialist laboratory can assist

Executive summary

THE HEADLINE POINT: GENAI READS AND CHARACTERISES AT SCALE UNDER INSTRUCTIONS YOU WRITE: THE INSTRUCTIONS ARE TESTED LIKE SEARCH TERMS, THE OUTPUTS VERIFIED LIKE A JUNIOR'S WORK, AND THE LAWYER REMAINS ACCOUNTABLE FOR ALL OF IT

Generative AI in disclosure is prompted review: a model reads documents against written instructions (the prompt) and returns assessments: relevant or not, privileged or not, a summary, a proposed redaction: at machine speed and lawyer-adjustable granularity. Three disciplines make it defensible. **Prompts are methodology:** the instruction text embodies the legal judgement, so it is drafted with the issues, tested against documents with known correct answers, refined on evidence and version-controlled: guide 84's term disciplines applied to instructions, per the ILTA guide's prompt-testing workflows. **Outputs are verified, not trusted:** AI assessments are sampled against human judgement at stated rates, disagreements diagnosed (prompt defect, document oddity, model failure), and the validation statistics of guides 85 and 89: precision, recall, elusion of the AI-negative population: carried over intact: the machine is a reader whose work is checked, never a decision-maker whose word is final. **Accountability never transfers:** Ayinde's line: lawyers answer for AI-generated content and face sanction for unverified reliance: plus the continuity-of-evidence and auditability expectations of PD 57AD mean the workflow documents what was prompted, what returned, what was verified and by whom. Run this way, GenAI is a proportionality instrument the DRD can recite; run as an oracle, it is a sanction waiting for its case name.

- **Prompted review, human-owned:** the model applies instructions lawyers wrote: the judgement stays in the drafting and the checking.
- **Test prompts like search terms:** benchmark sets, iteration on evidence, version control: the guide 84-85 method, transplanted.
- **Verify by sampling:** stated rates, both directions (AI-positive and AI-negative), with the guide 89 statistics carried over.
- **Document for audit:** prompts, versions, model identity, outputs, verification results: the continuity record the courts expect.
- **Data protection travels too:** disclosure data through AI systems raises UK GDPR and confidentiality questions answered before deployment, not after.

The problem in plain English: a reader that never tires and sometimes invents

A generative model does something no earlier disclosure technology did: it reads. Not ranks by resemblance (guide 89's TAR), not matches strings (guide 84's keywords): reads, in the working sense: give it a document and an instruction ("is this about the pricing arrangements? explain briefly") and it returns an assessment with reasons. Scaled across a corpus, that is a review team of unlimited size, perfect stamina and total consistency of mood: the proposition reshaping large-matter economics.

The same machinery has failure modes no earlier technology had either. It can be confidently wrong: fluent assessments resting on misreadings. It can be prompt-sensitive: the same document assessed differently under slightly different instructions. It can vary between runs. And its errors do not look like errors: they look like a careful junior's work, which is precisely what makes unverified reliance dangerous and why the Divisional Court's *Ayinde* line exists. The resolution is not abstinence and not faith: it is the same pattern this series has applied to every tool: written methodology, tested before use, verified by sampling, documented for audit. The novelty is the reader; the constitution is unchanged.

§ 0 3 · THE RULES OF THE ROAD

The guidance landscape: ILTA, the Law Society, the judiciary and Ayinde

- **The ILTA GenAI Best Practice Guide (September 2025):** the practitioner-drafted framework for GenAI in outgoing disclosure under PD 57AD: consensus best practices rather than rigid rules: covering prompt-testing workflows, continuity of evidence, and integration with active-learning review strategies: the reference document for §5-§7's disciplines, publicised through the Law Society's practical guidance of November 2025.
- **The accountability line: Ayinde (June 2025):** the Divisional Court in Ayinde v London Borough of Haringey (with Al-Haroun) confirmed that solicitors and barristers remain fully accountable for AI-generated content and face potential regulatory sanction for failing to verify accuracy: the principle that structures every workflow in this guide.
- **Judicial and professional guidance:** the judiciary's AI guidance (December 2023, updated March 2025) and the Law Society's generative AI essentials frame court-facing and client-facing use: hallucination risk named, verification duties emphasised, professional obligations under the SRA framework unchanged by the tooling.
- **PD 57AD's fit:** the practice direction's technology duty (reliable, efficient, cost-effective conduct of disclosure using technology) accommodates GenAI without naming it; its cooperation and auditability expectations: consistent, auditable, defensible workflows: are what the ILTA guide operationalises: and the Disclosure Review Working Group's late-2025 consultation on PD 57AD's operation gave AI a prominent place, so the rules layer remains in motion: checked per matter.
- **The standing consequence:** nothing in the landscape prohibits; everything in it conditions: use is expected to be disclosed, methodical, verified and documented: the same posture guide 89 §6 described for TAR, with the accountability line drawn more sharply still.

§ 0 4 · THE USE CASES

Where GenAI fits: relevance, privilege, summarisation and redaction support

- **Relevance assessment:** prompted first-pass review against issue-drafted instructions: as a triage layer beside or within CAL workflows (guide 89 §7's chain), as a second reader over human calls, or: on suitable matters and under the full §6 verification: as the primary pass with human review concentrated on AI-positive and borderline populations.
- **Privilege flagging:** prompted identification of likely-privileged material for specialist lawyer review: an accelerant into the guide 91 disciplines, never the determination itself: privilege calls remain lawyer judgements, per guide 89's line, underlined.
- **Summarisation and case intelligence:** document and thread summaries, chronology extraction, entity and relationship mapping: review-support product that speeds humans rather than replacing their calls: with summaries treated as leads to the documents, not substitutes for them.
- **Redaction assistance:** prompted identification of redactable content (personal data, out-of-scope confidential material) for human confirmation and guide 93's irreversibility engineering: the model proposes; the workflow disposes.
- **Fit by matter, honestly:** corpus size, issue complexity, data sensitivity and timetable decide where in this menu a matter sits: the suitability conversation of guide 89 §11, extended: including the answer "not here, not yet".

Prompt engineering as methodology: drafting, testing and versioning

- **Draft from the issues:** prompts written with the pleadings and the review guide in hand: definitions, inclusions, exclusions, output format: the legal judgement externalised into instructions, exactly as a review guide externalises it for humans.
- **Benchmark before deployment:** prompts run against a test set with known correct answers (human-coded documents spanning the difficulty range): agreement measured, disagreements read, the prompt refined: guide 85's test-before-agreement discipline, applied to instructions: the ILTA guide's prompt-testing workflow in practice.
- **Iterate on evidence, version on the record:** each prompt revision benchmarked and logged with its results: the prompt file as the guide 84 §7 protocol file's sibling: what was asked, in which version, of which documents, is reconstructable throughout.
- **Stability checks:** consistency measured across runs and document orderings; material instability met by prompt hardening, consensus runs or routing to human review: the non-determinism named and managed rather than discovered in cross-examination.
- **Model identity matters:** the model, version and configuration recorded per exercise: outputs are a function of instructions *and* engine: the guide 84 dialect rule, for readers.

Verification: sampling AI outputs against human judgement

- **Sample both directions:** AI-positive calls sampled for precision; AI-negative populations sampled for elusion: the guide 85/89 machinery unchanged: the model is a retrieval-and-assessment layer whose recall is estimated, never presumed.
- **Diagnose disagreements:** each sampled miss read for cause: prompt defect (fix and re-run), document oddity (route to humans), text-quality gap (guide 86's territory), genuine model failure (weight the reliance accordingly): misses as routing instructions, per the standing method.
- **Escalation tiers by consequence:** verification intensity scaled to what the output does: light for internal summaries, heavier for relevance determinations feeding production, heaviest for anything touching privilege or redaction: proportionality applied to checking itself.
- **Human decision points preserved:** production decisions, privilege determinations and redaction confirmations made by lawyers on documents (or verified populations), with the AI layer's role stated: the Ayinde-proof architecture: the model informs; identified humans decide.
- **The statistics in the certificate:** where AI-assisted review underpins disclosure, the certificate's foundation includes the verification numbers: agreement rates, sample designs, elusion results: guide 89 §5's recitals, with the reader swapped.

Transparency, data protection and the protocol

- **Disclose the methodology:** GenAI use declared in the DRD conversation like TAR before it: workflow, use cases, verification design: cooperation performed openly; silent deployment is the avoidable satellite dispute, per guide 89 §6's pattern.
- **Continuity of evidence:** the audit trail: prompts and versions, model identity, document populations processed, outputs, verification records, human decision points: maintained as first-class records per the ILTA guide's continuity workflows: guide 89 §8's record-keeping, extended to instructions and outputs.
- **UK GDPR and confidentiality:** disclosure data through AI systems is processing: lawful basis, minimisation, processor terms, security and residency of the AI platform, client confidentiality and any LPP implications of third-party processing: assessed and documented before deployment, with client instructions taken where the sensitivity warrants.
- **Cross-checking the opponent:** the §11 questions in reverse when their production was AI-assisted: use, workflow, verification, statistics: an exercise that cannot answer them meets guide 89 Example 2's fate with newer machinery.
- **Keep the rules current:** the landscape moves (the DRWG consultation, evolving guidance): per-matter verification of the current position, per this series' standing research rule.

Source architecture: where else the evidence lives

Per this series' source-architecture discipline, adapted to method: an AI-assisted exercise's defensibility, and its blind spots, live in mappable places. The defensibility lives in: the prompt file (versions, benchmarks, revision reasons); the verification records (samples, agreement rates, elusion results, disagreement diagnoses); the processing inventory (which populations went through which prompt under which model); the human decision log (who confirmed what); and the protocol and correspondence (what was disclosed and agreed). The blind spots live where every text tool's do, plus the reader's own: the unextracted and unOCRed layers (guides 81, 86: the model cannot read what processing never yielded); the non-text estate (media, structured data: guides 71-77's methods, stated in the methodology); the low-confidence and unstable populations (routed to humans by design); prompt scope itself (the instruction that never asked about the issue that later emerged: prompts re-run as issues evolve, like terms); and collection scope beneath everything (guide 89 Example 3's lesson: sometimes the corpus, not the reader, is the problem). Challenges and defences alike are archaeology across these records: which is why they are kept as evidence from the first prompt.

QUESTION	PROMPT FILE	VERIFICATION RECORDS	PROCESSING INVENTORY	HUMAN DECISION LOG	TEXT-QUALITY LAYER	SCOPE LAYERS
What was the model asked?	Y: versions + benchmarks	n/a	Prompt-to-population map	n/a	n/a	n/a
Was its work checked?	n/a	Y: the direct answer	Coverage confirmed	Confirmations recorded	n/a	n/a
What could it not read?	Instruction gaps	Elusion finds' character	Excluded populations	n/a	Y: guides 81/86 residuals	Non-text routes stated
Who decided?	n/a	n/a	n/a	Y: the Ayinde answer	n/a	n/a
Did the corpus, not the reader, miss it?	Prompt scope reviewed	Verification says reader; silence says scope	n/a	n/a	n/a	Y: collection guides govern

Fallback strategy in one line: keep prompts, verification and decisions as first-class records from day one: an AI-assisted exercise that can replay what it asked, what returned and who confirmed it survives scrutiny; one that cannot is Example 2 of guide 89 with a newer engine.

Worked examples

EXAMPLE 1 · THE PROMPTED FIRST PASS THAT SURVIVED ITS AUDIT

A 900,000-document commercial dispute ran GenAI as prompted first-pass relevance triage under a disclosed methodology: prompts drafted from the review guide, benchmarked against 1,200 human-coded documents until agreement stabilised, versions logged; the corpus processed with model identity recorded; AI-positive documents human-reviewed; the AI-negative population elusion-sampled per guide 85's arithmetic. The opponent's methodological challenge arrived mid-exercise and was answered from the running records in a single schedule: prompt versions with benchmark scores, verification rates in both directions, the elusion result, and the human decision architecture. The application was not pursued. The exercise closed at a review cost the linear estimate would have tripled, and the disclosure certificate recited the same schedule: the §8 records doing at corpus scale what they were kept for.

EXAMPLE 2 · THE SUMMARY THAT INVENTED A CONCESSION

A team preparing a witness statement relied on AI thread summaries of a 400-email negotiation chain. One summary reported the counterparty "conceding the delivery-date variation in the 14 March email": fluent, specific, and wrong: the 14 March email discussed a different contract, and no concession appeared anywhere in the chain. The error surfaced when the drafting lawyer applied the standing rule: summaries are leads, and every load-bearing assertion is verified against the documents: the guide 83 thread record showing the true sequence in minutes. The near-miss became the matter's training example: summarisation redeployed with citation-required prompts (every assertion linked to a message), a verification tier for anything entering work product, and the incident logged in the prompt file. The Ayinde principle, experienced rather than read: the lawyer who checked answered confidently; the lawyer who had not would have answered fluently.

EXAMPLE 3 · THE PRIVILEGE FLAGS THAT ACCELERATED, AND NEVER DECIDED

A regulatory disclosure with a hard deadline used prompted privilege flagging across 240,000 documents: instructions drafted with the privilege team (advice contexts, lawyer identities, the guide 91 categories), benchmarked on known-privileged and known-not sets, and applied as a routing layer: high-likelihood material queued first to specialist lawyers, the remainder to the standard stream with sampling. Every claim in the eventual privilege schedule was a lawyer's determination on a read document; the AI layer's role: disclosed in the methodology statement: was queue order and second-reader alerts, including catching four privileged threads a tired human pass had marked for production. The regulator's later question about AI use was answered in one paragraph precisely because the architecture had kept the machine on the right side of the line: accelerant, alert, never adjudicator.

Common mistakes and technical limitations

Common mistakes

- **Oracle deployment:** outputs relied on unverified: Example 2's invented concession, reaching a statement of truth.
- **Prompts untested:** instructions drafted in an afternoon and deployed corpus-wide: guide 84's Friday list, reborn as a prompt.
- **Verification one-directional:** AI-positives checked, the AI-negative population never elusion-sampled: recall presumed, guide 85 unlearned.
- **Silent use:** GenAI surfacing under challenge instead of in the DRD: the cooperation own-goal, newest edition.
- **Records reconstructed:** prompts and outputs assembled after the question arrives: §8's losing position.
- **Privilege delegated:** flags treated as determinations: the Example 3 line crossed, with everything that follows.
- **Data protection retrofitted:** client data through third-party AI platforms with the UK GDPR analysis done afterwards.
- **Summaries as evidence:** AI characterisations quoted in work product without document verification: the citation-required discipline absent.

Technical limitations

- **Confident error is structural:** fluency and accuracy are independent: verification exists because reading the output cannot reveal which you have.
- **Non-determinism is real:** runs vary: stability is measured and managed (§5), never assumed.
- **Context limits bound reading:** very long documents and sprawling families are read in windows: chunking strategies stated, thread-level questions routed to guide 83's structures.
- **The model reads what extraction yields:** OCR gaps, exceptions and non-text material bound it exactly as they bound TAR: guides 81/86/89's dependencies, unchanged.
- **Capability moves fast:** models, tools and guidance evolve on quarters, not years: per-matter currency checks over remembered positions, per the standing rule.

§ 11 · INTERROGATORIES & DRAFTING AIDS

Questions to ask · Suggested wording

Ask your client

- What is the data's sensitivity profile: and which AI platforms, terms and residencies can it lawfully and contractually touch?
- Which use cases does the matter actually need: triage, privilege acceleration, summarisation: and which carry consequences demanding the heaviest verification?
- Do any client policies, regulator expectations or counterparty agreements constrain AI use on this matter?

Ask your opponent

- Please state whether generative AI has been or will be used in your disclosure exercise, identifying the use cases, tools and models, and the workflow in which they operate.
- Please disclose the verification methodology applied to AI outputs, including sampling of populations assessed as non-responsive, the rates and results, and the human decision points in the workflow.
- Please confirm that records of prompts, versions, processed populations and verification results have been maintained and will be preserved.

Ask your eDiscovery / forensic provider

- What do the benchmark results show for the current prompt version: and what did the last revision fix?
- Are both verification directions running at the agreed rates: and what are the elusion finds saying?
- Where does the data go, under what terms: and does the §7 data-protection file cover it?

SUGGESTED WORDING · GENAI PARAGRAPH FOR THE DRD / PROTOCOL

"The parties may use generative AI tools in their disclosure workflows and shall disclose: the use cases (including relevance assessment, privilege identification, summarisation or redaction support); the tools and models used; the prompt development and testing methodology, with prompt versions recorded; and the verification methodology, including sampling of AI-assessed populations in both directions at stated rates and confidence, with results disclosed. Determinations of disclosability, privilege and redaction shall in every case be made or verified by identified lawyers, and the role of AI outputs in such determinations shall be recorded. Records of prompts and versions, models, processed populations, outputs and verification shall be maintained and preserved for the duration of the proceedings. Each party confirms that its use of such tools complies with its data protection and confidentiality obligations."

Checklist and red flags · When to involve a digital forensic expert

The GenAI checklist

- ☐ Use cases chosen by matter fit; consequence tiers mapped
- ☐ Data protection and confidentiality analysis completed pre-deployment
- ☐ Prompts drafted from the issues; benchmarked before corpus use
- ☐ Prompt versions, models and configurations logged throughout
- ☐ Stability measured; unstable populations routed to humans
- ☐ Verification sampling both directions at stated rates
- ☐ Disagreements diagnosed and routed, not just counted
- ☐ Human decision points preserved and recorded for privilege, production, redaction
- ☐ Methodology disclosed in the DRD conversation
- ☐ All records maintained as first-class evidence from day one

Red flags

- AI outputs in work product without document verification: Example 2's origin.
- Methodology statements that cannot name their prompt versions.
- No elusion sampling of AI-negative populations.
- Privilege schedules resting on flags rather than determinations.
- AI use surfacing for the first time under challenge.
- Client data on AI platforms nobody's data-protection file mentions.
- Opposing certificates over AI exercises with no verification numbers.

When to involve a digital forensic expert

At scoping, where use cases, platforms and the data-protection posture are engineered before deployment; at build, where prompts are benchmarked and the verification design installed (Example 1's architecture); through the exercise, where the records accumulate as evidence and the statistics are watched; and at challenge, in either direction, where AI-assisted methodologies are audited against their records and evidenced under CPR Part 35. Through **e-discovery.uk**, GenAI runs inside the same constitution as every tool before it: human judgement in charge, methodology disclosed, results verified: the machine's reading, defensible because the lawyers' checking was designed in.

§ 13 · COMMON QUESTIONS

Frequently asked questions

Is it actually permissible to use generative AI in court-ordered disclosure?

Yes: PD 57AD's technology duty accommodates it, the ILTA best-practice guide (September 2025) exists precisely to structure it, and the Law Society's practical guidance points practitioners to that framework: the conditions are methodology, verification, documentation and disclosure of use, not permission. What the authorities punish: Ayinde's line: is unverified reliance, not the tool.

How is this different from the TAR we already use?

TAR ranks by learned resemblance to coded examples; GenAI reads against written instructions and returns assessments with reasons: instructions you can draft, test and version like search terms: guide 89's constitution carries over, with the prompt file joining the training record. In practice they combine: ranking for triage, prompted reading for assessment, humans for decisions.

What is "hallucination" and how do we protect against it in disclosure work?

Confident output not grounded in the document: Example 2's invented concession is the disclosure-native form. Protections are structural: citation-required prompts, verification sampling at consequence-scaled rates, and the standing rule that summaries are leads: every load-bearing assertion checked against the source before it enters work product.

Do we have to tell the other side we used AI?

Under PD 57AD's cooperation model, effectively yes for outgoing disclosure: methodology transparency is the expectation the ILTA guide operationalises, and the DRD is the venue: what is disclosed is workflow and verification, not your prompts' work product or document-level outputs. Silent use discovered later is the most avoidable dispute in this guide.

Can AI make privilege calls?

It can flag, queue and alert: Example 3's architecture: and the determination remains a lawyer's on a read document: the line every current authority draws and the one place verification is never economised. An AI-accelerated privilege review is good practice; an AI-adjudicated one is a waiver argument waiting to happen.

Our client's data is highly sensitive. Does that rule GenAI out?

It rules the posture in: platform terms, residency, security and processor obligations assessed against UK GDPR and confidentiality duties before anything is processed: with deployment models ranging from restricted enterprise environments to not-at-all: §7's analysis decides, documented either way. Sensitivity shapes the how; it does not by itself answer the whether.

§ 1 4 · REFERENCE

Glossary

Benchmark set

Human-coded documents with known answers for prompt testing.

Citation-required prompting

Instructions demanding source references per assertion.

Consequence tier

Verification intensity scaled to an output's downstream effect.

Continuity of evidence

The preserved audit trail of prompts, outputs and decisions.

Elusion (AI-negative)

Responsive material in the population the model assessed as non-responsive.

Hallucination

Fluent output not grounded in the source document.

Prompt

The written instruction the model applies to documents.

Prompt file

The versioned record of instructions, benchmarks and revisions.

Prompted review

Document assessment by a model under drafted instructions.

Verification sampling

Human checking of AI outputs at stated rates, both directions.

Sources and authoritative references

The GenAI disciplines in this guide draw on materials published by our laboratory and e-discovery practice, referenced here as sources:

- e-discovery.uk, EDRM Phase 05 · Review (AI-assisted workflows, verification design and the record-keeping as practised): e-discovery.uk/edrm/review · Phase 04 · Processing: e-discovery.uk/edrm/processing
- cflab.uk, digital forensics services (methodology auditing, AI-exercise validation, CPR Part 35 evidence): cflab.uk/digital-forensics-services · guides library: cflab.uk/guides
- e-discovery.uk, practice method and Knowledge Centre: e-discovery.uk · e-discovery.uk/knowledge
- The Law Society, "Generative AI in legal disclosure: a practical guide" (November 2025): lawsociety.org.uk, [GenAI in legal disclosure](https://lawsociety.org.uk/genai-in-legal-disclosure) · "Generative AI: the essentials": lawsociety.org.uk, [GenAI essentials](https://lawsociety.org.uk/genai-essentials)
- ILTA, Generative AI Best Practice Guide (Outgoing Disclosure) (30 September 2025), and Courts and Tribunals Judiciary, AI guidance for judicial office holders (updated March 2025): judiciary.uk, [AI guidance](https://judiciary.uk/ai-guidance)
- PD 57AD: justice.gov.uk, [PD 57AD](https://justice.gov.uk/pd-57ad) · CPR Part 35: justice.gov.uk, [Part 35](https://justice.gov.uk/cpr-part-35) · ICO, UK GDPR guidance: ico.org.uk

DISCLAIMER

This publication provides general information about generative AI in disclosure as at August 2026. Tools, models, professional guidance and judicial expectations in this area develop rapidly, including through the Disclosure Review Working Group's ongoing consideration of PD 57AD; verify the current position for the matter at hand. The worked examples are fictional. This guide is not legal advice, does not create a professional relationship, and should not be relied upon in place of advice on the facts of a specific matter from a qualified lawyer or a suitably qualified forensic practitioner. Links were live at the date of publication.

§ HOW A SPECIALIST LABORATORY CAN ASSIST

Working with Computer Forensics Lab

GenAI at **e-discovery.uk** runs inside the constitution this series has applied to every technology: use cases scoped honestly (including "not here"); the data-protection posture engineered before a document moves; prompts drafted with the legal team, benchmarked, versioned and stability-checked; verification sampling both directions at consequence-scaled rates; human decision points preserved where the authorities require them (Example 3's architecture); and the whole record: prompts, models, populations, results, decisions: maintained as first-class evidence, so methodological questions are answered from files rather than memory (Example 1's schedule). The laboratory side audits opposing AI-assisted exercises and evidences methodology under CPR Part 35 where the machine's reading becomes the issue. Conflicts checks, confidential scoping discussions and fixed-fee estimates are routine; and if generative tools are already touching your live disclosure: or your opponent's: the verification architecture is this week's conversation, before the certificate depends on it.



I N S T R U C T T H E L A B

Speak to a forensic examiner, not a salesperson.

Describe your matter in confidence and an examiner will respond the same working day. Conflicts checks and scoping calls with US counsel are routine.

NEW ENQUIRIES

+44 (0)20 7164 6971

EMAIL

info@cflab.uk

E-DISCOVERY

e-discovery.uk

Computer Forensics Lab Ltd · Euro House, 133 Ballards Lane, Finchley, London N3 1LJ, United Kingdom

cflab.uk · Customer service +44 (0)20 7164 6915 · ICO ZB395023 · NPCC · ISO 17025 aligned practice · UK Gov Digital Marketplace