



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

Technology-Assisted Review for UK Lawyers

Predictive Coding, Continuous Active Learning and Defending the Machine's Judgement

Technology-assisted review turns human relevance decisions into a model that ranks the entire corpus: reviewers teach by example, the system learns what responsive looks like, and effort concentrates where responsiveness is likely. English courts approved predictive coding a decade ago, PD 57AD expects parties to consider it, and on large matters it is now the difference between a proportionate review and an impossible one. What courts have not licensed is the unexamined version: TAR without disclosed methodology, without validation, without the statistics that show the stopping decision was sound. This guide explains TAR for UK lawyers: how the workflows actually operate, where the human judgement sits, the validation arithmetic that makes results defensible (building on guide 85's disciplines), the negotiation and protocol questions, and how TAR coexists with keywords, threading and the rest of the toolchain.

⌚ 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**, runs technology-assisted review on large-corpus matters: workflow design, seed and training management, validation statistics and the protocol drafting that makes machine-ranked review defensible under PD 57AD: with the honest scoping advice about when TAR earns its complexity and when it does not.

ESTABLISHED 2007 · LONDON

ISO 17025-ALIGNED PROCEDURES

TAR & CAL WORKFLOWS

VALIDATION STATISTICS

CPR PART 35 EXPERT REPORTS

FULL CHAIN-OF-CUSTODY DOCUMENTATION

§ CONTENTS

In this guide

- 01 Executive summary
- 02 The problem in plain English: teaching a machine what relevant looks like
- 03 How TAR works: models, training and the two generations
- 04 The human layer: seeds, subject-matter reviewers and consistency
- 05 Validation: recall, precision and the stopping decision
- 06 TAR in the protocol: disclosure, negotiation and the courts
- 07 TAR in the toolchain: keywords, threading and privilege
- 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: TAR IS HUMAN JUDGEMENT AMPLIFIED BY RANKING: ITS DEFENSIBILITY IS THE TRAINING RECORD PLUS THE VALIDATION STATISTICS, AND BOTH ARE BUILT, NOT ASSUMED

TAR workflows share one architecture: senior reviewers code example documents; a model learns the statistical signature of responsiveness from text features; the corpus is scored and ranked; and review effort follows the ranking until validation shows the remainder can defensibly rest. Modern practice runs **continuous active learning (CAL)**: the model retrains as every batch is coded, always serving the most-likely-responsive documents next, tolerant of rolling data and shifting understanding: largely displacing the older fixed-training-set generation. Two things carry the weight. **The human layer**: the model is only as good as its teachers: seed and training decisions made by reviewers who know the case, coding consistency managed, borderline calls escalated: §4's disciplines. **The validation arithmetic**: the stopping decision: "the unreviewed remainder is acceptably unlikely to contain material documents": is proven by sampling, not asserted: elusion samples of the un-reviewed set, recall estimates with stated confidence, the guide 85 machinery at corpus scale: §5's statistics. English authority has approved the approach for a decade and PD 57AD expects it considered; what gets attacked is not TAR but undocumented TAR: methodology undisclosed, validation unrun, the stopping point a budget decision wearing a science costume. The protocol discloses the workflow, the record preserves the training history, and the certificate recites the numbers.

- **CAL is the current standard**: continuous retraining, richest documents first, robust to rolling collections: the generation to specify.
- **Teachers decide quality**: case-knowledgeable reviewers on training, consistency managed, disagreements resolved on the record.
- **Stopping is a statistical event**: recall estimated, elusion sampled, confidence stated: guide 85's disciplines are the licence to stop.
- **Disclose the methodology**: workflow, validation design and results shared per PD 57AD cooperation: secrecy is what loses TAR arguments.
- **TAR joins the toolchain**: after processing and dedup, beside keywords and threading, before privilege: §7's placement, not a replacement for everything.

The problem in plain English: teaching a machine what relevant looks like

A million-document review has a brutal property: most of it is irrelevant, and nobody knows which most. Linear review reads everything to find the responsive fraction: paying full price for every certain no. TAR inverts the economics: let experienced lawyers read a few thousand documents, let a model generalise their judgement across the million, and read in the order the model proposes: responsive material concentrating in the early batches, the tail thinning until what remains is statistically quiet.

Nothing mystical happens. The model learns which textual features co-occur with the lawyers' "responsive" calls: vocabulary, phrases, patterns: and scores unseen documents by resemblance. It knows nothing about the law, the issues or the truth; it knows what the teachers showed it. That is both the power and the whole risk register: taught well, validated honestly, it finds in weeks what linear review finds in quarters, at defensible recall; taught carelessly: skewed seeds, inconsistent coding, stopping when the budget ran out: it manufactures a confident-looking review with holes nobody measured. The difference between the two is not the algorithm: every section that follows is about the difference.

How TAR works: models, training and the two generations

- **The learning core:** documents represented as text features; a classifier trained on coded examples; every document scored for responsiveness likelihood: the ranking that drives everything else, re-computed as training grows.
- **TAR 1.0 (the first generation):** a fixed training set coded by seniors, the model applied once, documents above a cut-off reviewed: sensitive to seed quality and unstable on rolling data: encountered now mostly in older protocols and opposing methodologies worth questioning.
- **TAR 2.0 / CAL (the current standard):** continuous retraining with every coded batch, the queue always serving the highest-ranked unreviewed documents: no fixed training phase, graceful with rolling collections and evolving issue understanding, and the version this guide assumes throughout.
- **What the model sees and does not:** extracted text (guide 81's yield: the model inherits every OCR and exception gap per guides 86 and 81: stated in the methodology); not images, not audio, not the spreadsheet's formulas: the §7 toolchain covers what ranking cannot.
- **Scores are rankings, not verdicts:** a document's score orders the queue and supports the statistics: individual production decisions remain human ones on reviewed documents: the line that keeps the workflow inside every duty.

The human layer: seeds, subject-matter reviewers and consistency

- **Seeding without skew:** initial examples from known-responsive documents, targeted searches and random draws: diverse enough that the model learns the issues, not one document type: with the seed provenance recorded like guide 84's term origins.
- **Who codes training matters:** the model amplifies its teachers: case-knowledgeable reviewers on the training stream, the review guide written before coding starts, and the borderline categories (the grey documents that teach most) escalated rather than guessed.
- **Consistency management:** overturn rates tracked, disagreement sampled and resolved, the review guide amended on the record when issue understanding shifts: because inconsistent teaching is noise the model faithfully learns.
- **The audit trail:** who coded what, when, under which guide version: the training history preserved as the methodology's biography: what "disclosed workflow" means when someone asks two years later.
- **Issue drift handled openly:** pleadings amend, issues sharpen: CAL absorbs re-teaching, and the record shows the re-teaching: the flexibility is a feature precisely because it is documented.

Validation: recall, precision and the stopping decision

- **The question, precisely:** what fraction of the responsive material has the review found (recall), and what does the unreviewed remainder contain: answered by sampling, per guide 85's machinery at corpus scale.
- **The elusion sample:** a sized random draw from the unreviewed, low-ranked remainder, reviewed by humans: responsive finds counted, their character read (marginal stragglers or a missed vein?): the direct test of the stopping hypothesis.
- **Recall estimation:** from control sets or elusion arithmetic: the estimated found-fraction with stated confidence: the number the certificate leans on, reported as an estimate with its interval, never as false precision (guide 85 §10's honesty rule).
- **Stopping criteria agreed in advance:** the target recall and validation design set in the protocol before anyone is tired or over budget: the stopping decision executing a pre-agreed test, not rationalising an exhaustion.
- **Finds route, per the standing method:** elusion discoveries feed re-teaching or scope (a missed vocabulary, a channel the model never saw: guide 85 Example 2's logic): the cycle closing on evidence, and the record showing it closed.

TAR in the protocol: disclosure, negotiation and the courts

- **The authority in one line:** predictive coding has been approved in England and Wales since the mid-2010s, and PD 57AD's regime expects technology-assisted approaches to be considered on suitable matters: using TAR needs no permission slip; using it silently invites the fight.
- **What gets disclosed:** the workflow (CAL, the tool), the human-layer arrangements, the validation design and its results: methodology transparency, not document-level scores: the DRD's technology sections carrying it.
- **Negotiating the criteria:** target recall, elusion design and confidence levels agreed like search terms: numbers exchanged per guide 85's currency, disproportionate demands answered with review-cost arithmetic.
- **Receiving their TAR:** the same questions in reverse (§11's set): generation, teachers, validation, stopping numbers: an opposing methodology that cannot answer them is guide 85 Example 3's situation at corpus scale.
- **When challenges land:** attacks succeed against undocumented training, unrun validation and budget-shaped stopping: and fail against the record this guide builds: the pattern across the authorities, and the reason §4-§5 exist.

TAR in the toolchain: keywords, threading and privilege

- **After processing, on honest text:** TAR ranks what extraction yielded: the guide 81 pipeline and guide 86 text-quality work are upstream dependencies, and the methodology statement says what the model could not see.
- **With keywords, by design:** keywords for the nameable, TAR for the conceptual (guide 84 §5's division): protocols run them as complements: terms retrieving the known, ranking triaging everything, each covering the other's blindness.
- **With threading and dedup:** guide 82's suppression and guide 83's inclusiveness shape what the model trains on and reviewers see: settings aligned so the efficiencies stack instead of colliding.
- **Privilege stays human-led:** ranking can prioritise likely-privileged material for specialist review, but privilege calls and the guide 91-92 disciplines remain lawyer judgements: TAR accelerates the queue to them, never replaces them.
- **Beyond classic TAR:** newer AI-assisted review approaches (including generative tools: guide 90's subject) join the chain under the same constitution: human judgement, disclosed methodology, measured validation: the constants that outlive any particular technology.

Source architecture: where else the evidence lives

Per this series' source-architecture discipline, adapted to method: a TAR workflow's defensibility, and its gaps, both live in identifiable places. The defensibility lives in: the training record (seeds, coding history, guide versions: §4's biography); the validation file (sample designs, draws, results, the stopping decision's arithmetic); the protocol and correspondence (what was disclosed and agreed when); and the platform's own logs (rankings over time, batch composition: the workflow's black box recorder). The gaps live where the model could not see: the unextracted and poorly OCR'd text (guides 81 and 86's residuals: ranked blind or not at all); the non-text estate (media per guide 77, structured data per guides 71-74: routed to their own methods and said so); the low-ranked remainder itself (bounded by elusion, never proven empty); and the sources the corpus never included (collection scope, the oldest question, sitting under every review method equally). Challenging or defending a TAR exercise is archaeology across exactly these locations: which is why they are maintained as records rather than reconstructed as memories.

QUESTION	TRAINING RECORD	VALIDATION FILE	PROTOCOL & CORRESPONDENCE	PLATFORM LOGS	TEXT-QUALITY LAYER	NON-TEXT ROUTES
Was the model taught properly?	Y: the direct answer	Consistency metrics	Review guide versions	Coding chronology	n/a	n/a
Was stopping sound?	n/a	Y: recall + elusion arithmetic	Pre-agreed criteria	Ranking curves	Blind-spot bounds stated	n/a
What could the model not see?	n/a	Elusion finds' character	Methodology statement	Exception links	Y: guides 81/86 residuals	Y: guides 71-77 methods
Is their TAR defensible?	Demand it: §11	Demand the numbers	What they disclosed	Discoverable on challenge	Demand their posture	Demand their routes
Did scope, not ranking, miss it?	n/a	Elusion says ranking; silence says scope	The DRD's source list	n/a	n/a	Collection guides govern

Fallback strategy in one line: keep the training and validation records as first-class evidence from day one: a TAR exercise that must reconstruct its own history is already losing the argument about it.

Worked examples

EXAMPLE 1 · THE 1.8 MILLION DOCUMENTS REVIEWED BY 140,000

A competition follow-on claim faced 1.8m documents post-dedup and a review estimate its budget could not survive linearly. The CAL workflow of this guide ran instead: a two-partner training stream with a written review guide, seeds from the pleadings' key documents plus random draws, retraining per batch. Responsiveness concentrated exactly as the method predicts: the first 90,000 ranked documents held most of what mattered; by 140,000, batches ran overwhelmingly non-responsive. The pre-agreed stopping test then executed: an elusion sample of the remainder at stated confidence, finding a residual rate consistent with the target recall and the finds all marginal. The disclosure certificate recited the workflow, the numbers and the elusion result; the opponent's methodological questions were answered from the running records in one letter; and the review closed at under a tenth of the linear estimate with statistics the certificate could stand on.

EXAMPLE 2 · THE STOPPING DECISION THAT WAS A BUDGET IN COSTUME

An opposing party disclosed that TAR had been used and review "concluded when the tool indicated completion". The §11 questions unpicked it: no pre-agreed recall target existed; the "indication" was a ranking-curve flattening read by the provider; no elusion sample of the unreviewed 700,000 documents had ever been drawn; and the training stream had passed between three agencies mid-exercise with no consistency reconciliation. The court ordered a validation exercise: an independent elusion sample, which found responsive material at a rate incompatible with any respectable recall, concentrated in a vocabulary the late training had drifted away from. Re-teaching and a further review tranche followed, at the producing party's cost. The judgment's language was the §6 pattern exactly: the criticism was never that a machine was used, but that nobody could show what it had done.

EXAMPLE 3 · THE ELUSION FIND THAT WAS A SCOPE HOLE

A claimant's CAL exercise validated cleanly on recall but its elusion sample contained two responsive documents of an odd character: both were emails *about* attachments: engineering change notices: whose attachments were absent, and both referenced a numbering series almost unrepresented in the corpus. The §8 routing read the signal correctly: not a ranking failure (the emails had scored low because their bodies were administrative) but a scope question. Investigation found the change-notice repository: a legacy engineering system outside the collected sources: guide 74's territory, never mapped at the DRD stage. The collection was extended, 9,000 notices entered a supplementary exercise, and the eventual certificate recited both the original validation and the scope correction. The example's moral is §5's last bullet: elusion finds are read for what they say, and sometimes they say the corpus, not the model, was the problem.

Common mistakes and technical limitations

Common mistakes

- **Stopping by budget or curve-gazing:** Example 2's costume: no pre-agreed test, no elusion, no numbers.
- **Junior or fragmented training streams:** the model taught by whoever was available: consistency never managed.
- **Methodology kept secret:** TAR run silently and surfaced under challenge: the cooperation own-goal at its most expensive.
- **Validation theatre:** tiny convenience samples dressed as statistics: guide 85 §10's sorted-browse "samples", promoted to corpus scale.
- **Text-quality gaps unstated:** the model ranking blind over unOCRed archives while the methodology statement claims the corpus: guides 81/86 unjoined.
- **Elusion finds unrouted:** Example 3's signal ignored: discoveries logged and never read for scope.
- **TAR asked to do privilege:** ranking treated as a privilege call: the human-led line of §7 crossed.
- **Records reconstructed, not kept:** the training biography assembled after the challenge arrives: §8's losing position.

Technical limitations

- **The model reads text only:** images, audio, formulas and structure are invisible: the toolchain covers them and the statement says so.
- **Rare document types train poorly:** a handful of examples of a critical class may never rank well: targeted searches guard the corners keywords were built for.
- **Recall estimates carry intervals:** the statistics bound, they do not guarantee: certificates state estimates as estimates.
- **Low prevalence strains everything:** when responsive material is vanishingly rare, validation sample sizes balloon: the design anticipates it or the confidence quietly evaporates.
- **Tools differ under one label:** "TAR" spans implementations with different behaviours: the tool and version belong in the protocol, per the standing dialect rule.

§ 11 · INTERROGATORIES & DRAFTING AIDS

Questions to ask · Suggested wording

Ask your client

- Is the corpus TAR-shaped: volume, text-heavy, prevalence: or would keywords and targeted review serve at this scale?
- Who are our teachers: which case-knowledgeable reviewers can staff the training stream consistently?
- What recall target and validation confidence does the matter's risk actually warrant: before the protocol fixes them?

Ask your opponent

- Please state whether technology-assisted review has been or will be used, identifying the tool, the workflow generation (including whether continuous active learning), and the arrangements for training and coding consistency.
- Please disclose the validation methodology and results, including target recall, sample designs, sizes, confidence levels, elusion results and the criteria by which review was or will be concluded.
- Please state what categories of material were not amenable to the model (including image-only, non-text and excepted items) and the methodology applied to each.

Ask your eDiscovery / forensic provider

- Are the training and validation records being kept as first-class evidence: who coded what, under which guide version?
- What are today's ranking curve, prevalence estimate and projected stopping numbers: and does the pre-agreed test still fit?
- What did the elusion finds actually say: marginal stragglers, a teaching gap, or Example 3's scope hole?

SUGGESTED WORDING · TAR PARAGRAPH FOR THE DRD / PROTOCOL

"The parties may use technology-assisted review, and shall disclose: the tool and workflow (including whether continuous active learning is used); the arrangements for training, including the seniority of coding reviewers and the management of coding consistency; and the validation methodology, comprising a target recall of [X]% at [Y]% confidence, elusion sampling of the unreviewed population at a size sufficient for that confidence, and disclosure of the resulting estimates. Review shall not be concluded except upon satisfaction of the agreed validation criteria, and responsive documents identified in validation sampling shall be addressed by further training, review or scope investigation as their character requires. Categories of documents not amenable to text-based ranking shall be identified and addressed by stated alternative methodologies. Training and validation records shall be maintained and preserved for the duration of the proceedings."

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

The TAR checklist

- ☐ Suitability assessed honestly: volume, prevalence, text share
- ☐ CAL workflow and tool specified in the protocol
- ☐ Training stream staffed by case-knowledgeable reviewers under a written guide
- ☐ Seed provenance and coding history recorded from day one
- ☐ Consistency tracked; disagreements resolved on the record
- ☐ Stopping criteria pre-agreed: recall target, confidence, elusion design
- ☐ Text-quality and non-text gaps stated with their covering methods
- ☐ Elusion finds routed: re-teach, re-review, or scope
- ☐ Methodology and results disclosed per the protocol
- ☐ Training and validation records preserved as evidence

Red flags

- Review "concluded when the tool indicated": Example 2's phrase.
- No elusion sample anywhere in a methodology statement.
- Training streams that changed hands without reconciliation.
- Recall claimed without confidence intervals or sample sizes.
- Methodology statements silent on OCR posture and exceptions.
- Elusion finds logged but never characterised: Example 3's signal, missed.
- TAR disclosed for the first time under challenge.

When to involve a digital forensic expert

At suitability and protocol stage, where the workflow, targets and validation design are engineered before commitments (Example 1's pre-agreed test); through the exercise, where training discipline and the running records are maintained and the statistics watched; at validation, where stopping executes its test and finds are routed (Example 3's reading); and at challenge, in either direction, where methodologies are audited against their records and evidenced under CPR Part 35 (Example 2's exercise). Through **e-discovery.uk**, TAR runs with the record-keeping as standard: the machine's judgement, defensible because the humans' judgement around it was documented.

§ 1 3 · COMMON QUESTIONS

Frequently asked questions

Is predictive coding actually approved by the English courts?

Yes: approval dates from the mid-2010s authorities and the current PD 57AD regime expects technology-assisted approaches to be considered on suitable matters: the live questions are always methodology and validation, never permission. Parties lose TAR arguments by running it secretly or stopping it unscientifically, not by using it.

How big does a matter need to be before TAR makes sense?

There is no statutory floor: the judgement weighs volume against prevalence, text share and timetable: as a working shape, six figures of post-dedup documents usually rewards CAL, low five figures usually does not, and between them the review-cost arithmetic decides. An honest provider will sometimes answer "keywords and targeted review", and that answer is worth having.

Can the model miss the smoking gun?

It can rank oddly anything unlike its training: which is why the method never relies on ranking alone: keyword retrieval guards the nameable, targeted searches guard rare critical classes, elusion sampling measures the remainder, and Example 3 shows the finds being read for scope. The linear-review comparator misses things too: it just never measures what.

Do we have to tell the other side we are using TAR?

Under PD 57AD's cooperation model, effectively yes: methodology transparency is the expectation, the DRD provides the fields, and undisclosed TAR surfacing later is the single most avoidable way to convert a sound review into a satellite dispute. What is disclosed is workflow and validation: not your document scores or coding decisions.

What recall target should we agree?

Commonly agreed targets sit in the range where the marginal documents are overwhelmingly duplicative or peripheral, with the precise figure and confidence negotiated against the matter's stakes and prevalence: the honest framing is that the target prices residual risk, and guide 85's arithmetic makes the pricing explicit. Refuse false precision: an estimate with a stated interval beats a confident bare percentage.

Is TAR the same as using generative AI for review?

No: classic TAR ranks by learned classification and is decade-settled law; generative approaches read and characterise documents and arrive under newer guidance with their own verification disciplines: guide 90 treats them. What transfers unchanged is the constitution: human judgement in charge, methodology disclosed, results validated by sampling: the parts of this guide that will outlive every tool it mentions.

§ 1 4 · REFERENCE

Glossary

CAL

Continuous active learning; retrain-per-batch TAR, the current standard.

Control set

A coded random sample used to measure model performance.

Elusion

Responsive material in the unreviewed remainder; sampled at stopping.

Prevalence

The responsive fraction of the corpus; drives sample sizes.

Ranking

The scored ordering of documents by responsiveness likelihood.

Recall

The fraction of responsive material found; the stopping statistic.

Review guide

The written coding standard the teachers apply.

Seed set

The initial examples that begin training.

Stopping criteria

The pre-agreed statistical test that concludes review.

Training stream

The human coding that teaches the model.

Sources and authoritative references

The TAR 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 (TAR and CAL workflows, validation and stopping as practised): e-discovery.uk/edrm/review · Phase 04 · Processing (the text the model sees): e-discovery.uk/edrm/processing
- cflab.uk, digital forensics services (methodology auditing, validation exercises, 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
- PD 57AD and the Disclosure Review Document (technology and methodology sections): justice.gov.uk, [PD 57AD](https://justice.gov.uk/pd-57ad)
- *Pyrrho Investments Ltd v MWB Property Ltd* [2016] EWHC 256 (Ch) and *Brown v BCA Trading Ltd* [2016] EWHC 1464 (Ch) (predictive coding approved): bailii.org, [Pyrrho](https://bailii.org/pyrrho) · bailii.org, [Brown](https://bailii.org/brown)
- CPR Part 35: justice.gov.uk, [Part 35](https://justice.gov.uk/part-35) · ICO, UK GDPR guidance: ico.org.uk

DISCLAIMER

This publication provides general information about technology-assisted review as at August 2026. Tool behaviour, workflow terminology and judicial guidance develop continuously; verify the current position for the platform and authorities in 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

TAR at **e-discovery.uk** runs to this guide's constitution: honest suitability advice first (including "not TAR" where the arithmetic says so); CAL workflows with case-knowledgeable training streams, written review guides and consistency management; stopping criteria agreed in the protocol and executed as statistics, with elusion finds read for what they mean (Example 3's routing); and the training and validation records kept as first-class evidence from the first batch: Example 1's one-letter answer to methodological questions as the standard deliverable. The laboratory side audits opposing exercises against their records: or their absence: and evidences methodology under CPR Part 35 where the machine's judgement becomes the issue, Example 2's exercise in either chair. Conflicts checks, confidential scoping discussions and fixed-fee estimates are routine; and if a seven-figure document population has just landed on your matter, the suitability conversation is this week's call: before the linear review estimate becomes anyone's anchor.



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