

PAPER TWO · MEDIA AND FILM PRODUCTION

Responsible AI Adoption in Media and Film Production

Why the obligations attach to the asset, and why the bill arrives at delivery

Paul Forrest

Founder, Ecaveo
September 2026

About this paper

This is the second paper in a series on the responsible deployment of artificial intelligence in regulated and evidence-led organisations. Each paper stands alone. Together they set out a single argument, which is that the durable advantage from artificial intelligence comes from better decisions, and that better decisions require systems designed around the person who is accountable for them.

Screen production follows population health research because it is the sector where the gap between what is claimed and what is evidenced is widest, and where the consequences of an undocumented decision surface furthest from the person who made it.

Editorial balance. Roughly eight parts in ten is evidence and sector analysis, and two parts in ten is original framework. Readers who want the argument without the operating model can stop at Chapter 10.

HOW TO READ THE EVIDENCE TAGS

Every substantive claim carries a tag stating what kind of evidence supports it. The tags run across the series, adapted to each sector.

PEER-REVIEWED A primary study in a peer-reviewed journal or conference.

REVIEW A systematic review, a meta-analysis or a commissioned sector study.

SURVEY Survey or polling data, cited with its sample and fieldwork dates.

LEGISLATION An Act, a regulation, a statutory instrument or a court judgment.

REGULATOR A published position from a regulator, funder or public body.

AGREEMENT A collective agreement, a union position or an industry contract term.

STANDARD A technical specification from a standards body or industry consortium.

PRACTICE Practitioner opinion, including the author's own.

ECAVEO The author's own framework, proposed and not yet proven.

A claim resting on two kinds of evidence carries both tags. Where the evidence is thin the paper says so, and here it is thin exactly where the spending is heaviest.

Legal, regulatory and contractual notice. This paper describes copyright law and policy, the EU AI Act, data protection law, collective agreements and litigation, in particular in Chapters 6, 7 and 9. Nothing in this document constitutes legal advice. The provision of legal advice sits outside the terms of any engagement with the author or with Ecaveo, and any question of a particular production's compliance, rights position or contractual obligations is a matter for its own qualified advisers and business affairs team.

Statements about what a source says are attributed to that source. Where the author draws a practical implication the source does not itself state, the text says so. Case law described here is unsettled, some of it is under appeal, and the position is the one the author was able to verify on 11 September 2026.

ILLUSTRATIVE MATERIAL

Every production, company and scenario invented for this paper is marked **illustrative scenario** where it appears. Where a real production is discussed, it is discussed only on the basis of public statements by the filmmakers or the companies involved, and those sources are cited. Published research is cited and numbered, and invented material is never presented as evidence.

CITATION

Forrest, P. (2026). *Responsible AI Adoption in Media and Film Production: Why the Obligations Attach to the Asset, and Why the Bill Arrives at Delivery*. Ecaveo AI Adoption Whitepaper Series, Paper Two. Version 1.0, September 2026.

Comments and corrections are welcome and will be acknowledged in later versions.

ALSO IN THIS SERIES

Paper one. *Responsible AI Adoption in Population-Scale Health Research*. Why evaluation capability decides how far adoption can go.

In preparation. *The Governed Agent*. Bounded autonomy in regulated operations.

In preparation. *Renting the Model, Owning the Learning*. Where advantage actually accumulates.

Contents

FRONT MATTER

A note from the author	4
Executive summary	6

PART ONE · WHAT THE SECTOR IS ACTUALLY BUYING

1	The clearance cliff	8
2	The adoption rate that does not exist	10
3	Where the evidence actually is	13
4	Where the evidence is absent	16
5	The jagged frontier	18

PART TWO · RIGHTS, PEOPLE AND AUDIENCE

6	Copyright and training data, unsettled everywhere	21
7	Performers, likeness and the scope of consent	23
8	The workforce, and which crafts are exposed	26
9	Disclosure, and why the flat label backfires	28
10	Provenance, and what marking cannot do	31

PART THREE · BUILDING THE CAPABILITY

11	The asset manifest	34
12	An operating model	36
13	Evaluation and assurance	39
14	A portfolio built to the evidence	41
15	Delivery over six to nine months	44
16	The intelligent production programme	46

END MATTER

Board and executive checklist	48
Glossary	49
Method and evidence grading	50
About the author	52
References	53

A NOTE FROM THE AUTHOR

The bill arrives at delivery

I came to this paper expecting to write about generative video, because that is where the sector's anxiety sits and where the money is being raised. The evidence sent me somewhere else entirely.

There is no peer-reviewed study with a baseline and a comparison for artificial intelligence in visual effects, rotoscoping, restoration, previsualisation or virtual production. Not a weak one. None. Every productivity figure circulating in the trade press for those crafts originates with a company selling the tools. **PRACTICE** Meanwhile the strongest evidence in the sector sits in subtitling, transcription and voice, which are the crafts with the least protection and the fewest advocates.^{9,10,11} **PEER-REVIEWED**

The second surprise was the disclosure research. A preregistered study of 4,976 people found that labelling content as AI-generated reduced perceived accuracy by a similar amount whether the content was true or false, and that human-made work mislabelled as AI suffered the same penalty.¹⁵ **PEER-REVIEWED** A flat label is not a neutral act of honesty. It is a design decision with a measurable cost, and the cost falls on truthful work.

The third was legal. Two judges in the same district decided two days apart that acquiring training data from pirated sources was, and was not, severable from transformative use.^{20,21} **LEGISLATION** In the United Kingdom the government abandoned its preferred copyright option in March 2026 and left the status quo in place.¹⁷ **REGULATOR** Nobody making a film today can rely on a settled answer about the provenance of the models they are using.

What connects all of this is timing. A decision taken in a development meeting or a post-production suite is paid for months later, at quality control, at clearance, at insurance, at platform delivery, at territory release or during an awards campaign. The people who take the decision are rarely the people who receive the bill, and by the time it arrives the evidence needed to answer it has usually been lost.

TWO PAGES THAT CIRCULATE ON THEIR OWN

Executive summary

Screen production has a governance problem that a corporate AI policy cannot solve, because the obligations in this sector do not attach to the organisation. They attach to the asset. A shot, a voice, a face, a cue, a subtitle track and a cut each carry their own rights, consents, disclosures and territorial restrictions, and each can be lawful in one use and undeliverable in another.

This paper argues that the capability worth building first is the ability to say what a frame is made of. It argues equally firmly against several things the sector is currently being sold, and it says plainly where the evidence for the sector's own enthusiasm does not exist.

1

There is no evidence where the spending is. No peer-reviewed study with a baseline and comparison exists for AI in visual effects, rotoscoping, restoration, previsualisation or virtual production. Every figure in circulation comes from a vendor.

2

The adoption rate is an artefact of the question. UK figures run from under 2 per cent of BFI fund applicants declaring AI use in the funded work to 57 per cent of Pact members reporting it and over 70 per cent in a trade survey that did not separate assistive from generative use.

3

A flat disclosure label damages truthful work. In a preregistered study of 4,976 people, an AI-generated label cut perceived accuracy by a similar amount for true and false headlines alike, and human-made work mislabelled as AI took the same penalty.

4

Provenance marking is a convenience, not a control. In the regulator's own testing, cropping an image removed the watermark. Peer-reviewed work shows invisible watermarks are provably removable while preserving image quality.

5

The United Kingdom has no bargained AI standard. SAG-AFTRA's 2026 agreement set a floor covering scanning, synthetics and biometric data. The equivalent UK negotiation remains unresolved, with a 99.6 per cent indicative ballot behind the performers' union.

6

The sustainability standard cannot see the problem. The screen sector's principal carbon calculator contains no treatment of AI, cloud computing or data centre emissions, while image generation uses over 60 times the energy of text generation.

What follows for a board, a funder or a programme

- **Record what a frame is made of, at the moment it is made.** Sources, rights basis, consents, the tool and version, the transformation applied, the reviewer and the permitted uses. None of this can be reconstructed later, and all of it is asked for at delivery.
- **Tier by clearance exposure.** The question to ask is what has to be true at delivery, in every territory the work will reach, for the asset to be usable. Whether a model was involved is a separate question and a less useful one.
- **Disclose the function, not the fact.** The evidence supports saying what the system did. It does not support a bare statement that artificial intelligence was used, which depresses trust in truthful work.

- **Refuse four things.** No performer scan without a bounded, specific and revocable consent. No model treated as a rights clearance. No generative element in a delivered asset without a recorded provenance entry. And no reliance on watermarking as an enforcement mechanism.

Chapter 16 proposes a production intelligence design workshop for a studio, a production company, a funder or a post facility. It begins with a single asset and asks what can be proved about it, which is unglamorous, cheap, and the part everybody skips.

IF YOU HAVE TEN MINUTES

Read Chapter 1 for the problem, Figure 8 for the record that answers it, and the board and executive checklist at the back. Those three between them carry the argument.

IF YOU COMMISSION OR FUND

Chapters 2, 6 and 9 are the ones that bear on what you can require of an applicant, and Chapter 14 sets out which uses have evidence behind them and which are being bought on a vendor's word.

IF YOU RUN A FACILITY OR A SHOW

Chapters 11 to 13 are the operational core. Chapter 15 says what to build first and Chapter 5 says why a bounded pilot beats a departmental rollout.



PART ONE

What the sector is actually buying

Five chapters on the distance between what is claimed for artificial intelligence in screen production and what has actually been measured.

CHAPTER ONE

The clearance cliff

A decision taken in a development meeting is paid for at delivery, by people who did not take it and cannot reconstruct it.

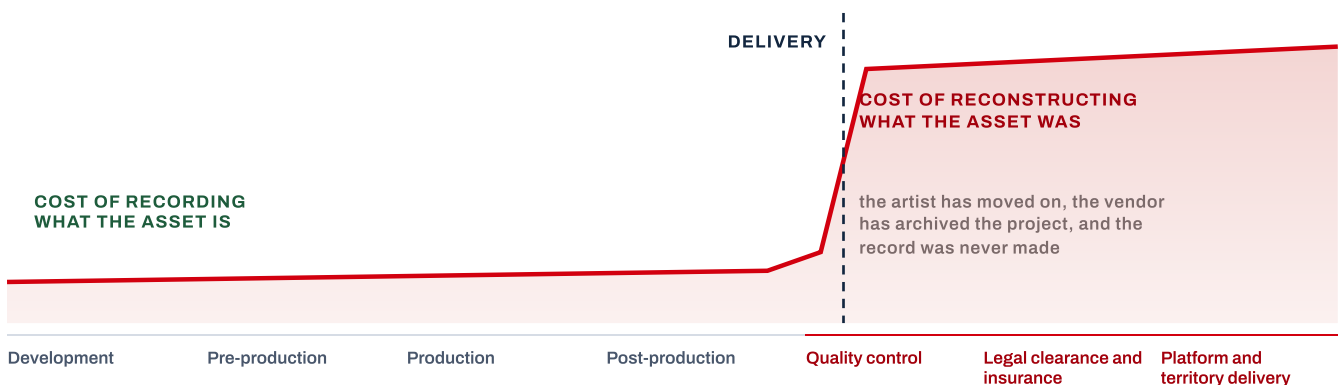
An artificial intelligence decision in screen production is cheap at the moment it is taken and expensive at the moment it is examined, and those two moments are months apart, with several companies and a change of personnel in between. A generative element enters a shot in a post-production suite on a Tuesday. The question of whether it can be delivered arises at quality control, at legal clearance, at errors and omissions insurance, at platform delivery, at territorial release or, most uncomfortably, during an awards campaign. By then the artist has moved to another show, the vendor has archived the project and the evidence that would answer the question was never recorded.

This is the central operational problem of the sector, and it is not a problem about models. It is a problem about when cost is incurred relative to when it is discovered. I call it the clearance cliff, and the argument of this paper is that the capability worth building first is the one that flattens it, which is the ability to say what a frame is made of at the moment the frame is made. **ECAVEO**

WHEN THE COST IS INCURRED, AND WHEN IT IS DISCOVERED

WHERE THE DECISION IS TAKEN

WHERE THE QUESTION IS ASKED



A minute of work, at the moment of creation

Every field the manifest of Chapter 11 asks for is held by somebody in the room while the asset is being made.

Days of work, or an unusable asset

Reconstruction runs against a release date, and where it fails the sequence is replaced, re-shot or cut.

FIGURE 1 The gap between the two is where undocumented AI use becomes expensive. Every gate on the right asks a question that can only be answered with evidence captured on the left.

Ecaveo framework. **ECAVEO**

1.1 Why a corporate policy will not hold

Most organisations respond to this by writing an artificial intelligence policy. In screen production a corporate policy is the wrong instrument, because the obligations in this sector do not attach to the organisation. They attach to the asset.

A single generated frame or cloned voice may implicate the rights position of the material a model was trained on, a performer's contractual and biometric rights, the confidentiality of unreleased production material, a downstream licensing term, a territorial disclosure rule and an audience's expectation of authenticity. Those obligations belong to that frame and that voice. They travel with the asset through vendors, versions, territories and windows, and they outlive the company that made it. A policy stating that the organisation uses artificial intelligence responsibly answers none of the questions that will actually be asked about that frame.

The same asset is lawful and unlawful at once. A sequence cleared for a domestic theatrical release may fail a platform's delivery specification, may require disclosure in one territory and not another, and may be unusable in a marketing cut because the consent that covered the film did not cover advertising. Nothing about the asset changed. The question changed, and the answer was never written down.

1.2 What this paper argues

Three propositions follow, and Parts One and Two earn them. The first is that the sector's confidence is running well ahead of its evidence, and in a specific and correctable direction. The second is that the obligations which matter most attach to people and to assets. Systems are where the sector has put its governance effort so far, which is why consent scope and provenance are the two disciplines worth building before anything else. The third is that disclosure, which is the sector's instinctive remedy, damages truthful work when it is designed as a flat label and helps when it is designed as a statement of function. **ECAVEO**

THE ARGUMENT IN ONE PARAGRAPH

Screen production has claim abundance and evidence scarcity, and its obligations attach to assets rather than to organisations. The useful contribution of artificial intelligence governance is therefore a record. Capture sources, rights basis, consent scope, tool and version, transformation and reviewer at the moment an asset is made, then tier by what has to be true at delivery. Disclose what the system did. Everything in this paper is either evidence about how good these tools actually are, or an argument about what a production should refuse to do with them.

CHAPTER TWO

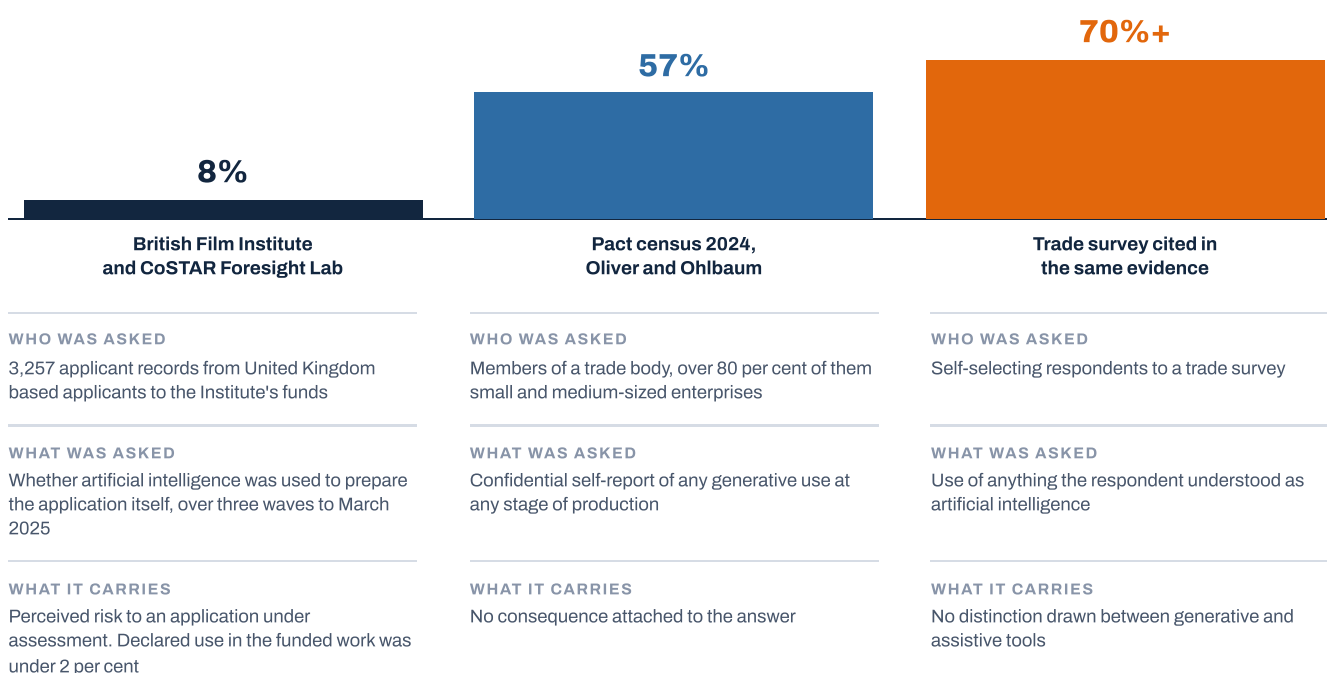
The adoption rate that does not exist

Published UK figures run from under 2 per cent to over 70 per cent, and they do not measure the same thing. The gap between what is declared and what is used is the governance problem itself.

Any executive asking how widely artificial intelligence is used in the United Kingdom screen sector will be given a number, and the number will be wrong in a way that depends entirely on who was asked and how.

Three figures are in circulation, and each of them is accurate. The British Film Institute, working with the CoSTAR Foresight Lab, reported that 8 per cent of United Kingdom based applicants to its funds declared using artificial intelligence to prepare their applications, rising from 4.6 per cent to 10.1 per cent across three reporting waves, while fewer than 2 per cent declared any intended use of artificial intelligence in the funded work itself, drawn from an analysis of 3,257 applicant records between August 2023 and March 2025.² **REVIEW** The Pact census for 2024, conducted by Oliver and Ohlbaum and submitted in written evidence to Parliament, found that 57 per cent of members, of which over 80 per cent are small and medium-sized enterprises, had used generative artificial intelligence at some point during the production process.³ **SURVEY** A trade survey cited in the same evidence put the figure above 70 per cent, and Pact itself notes that it did not distinguish between generative and assistive artificial intelligence.³ **SURVEY**

THREE UK ADOPTION FIGURES, AND WHAT EACH ONE MEASURES

**The three figures do not disagree about the world**

They measure three different things, and the first does not measure use in the work at all. The distance between what is declared and what is reported confidentially is the undeclared middle, which is the part of the sector using these tools with no route to say so.

FIGURE 2 Question wording, the population sampled and perceived risk produce the spread, and nobody disagrees about the facts. The first two columns do not even measure the same thing, which is the point. The distance between what is declared on a funding form and what is reported in a confidential census is the undeclared middle, and that is where governance is absent.

References 2 and 3. [REVIEW](#)

The spread is not disagreement about the world. The British Film Institute figure counts formal declarations about how an application was prepared, which carry perceived risk and cover a specific population applying for public money. Its figure for use in the funded work itself is under 2 per cent. The Pact figure counts self-reported use of any kind, at any stage, in a confidential census. The trade figure counts anything a respondent thought of as artificial intelligence, including features embedded in software they already owned. Anyone citing a single adoption number for this sector is citing an artefact of question wording.

PRACTICE

The operational reading matters more than the statistical one. Adoption is high and rising in informal and undeclared use, and low in formal declaration, and the distance between the two is precisely the population that has no route to record what it did. A producer who used a generative tool in development and did not declare it has not behaved badly. In most cases nobody asked, no form had a field for it, and the funding guidance

THE UNDECLARED MIDDLE

The interesting population sits between the fewer than 2 per cent who declare use in the work and the 57 per cent who report it confidentially. Most of that gap is people with no route to say what they did.

that now requires disclosure post-dates the work.

That guidance exists. The British Film Institute's guidance on the use of artificial intelligence in funding applications and funded projects, last updated in January 2026, provides that the Institute does not prohibit the use of artificial intelligence in applications or funded projects, but requires applicants to be transparent about the use of any artificial intelligence, both assistive and generative, and to operate within the law. It states that all funding applications contain questions on the use of artificial intelligence, and that awarded applicants are expected to make the Institute aware of its usage throughout the life of the project.²⁶

REGULATOR The requirement runs for the life of the project, and most production paperwork is not built to support that.



Pre-production is where the declared and the undeclared diverge most sharply. It is also the stage at which a record costs least to create.

CHAPTER THREE

Where the evidence actually is

The crafts with the strongest measured evidence are subtitling, transcription and voice, which are also the crafts with the least protection.

Three screen workflows have been measured properly. All three sit in the parts of the pipeline that attract the least attention and carry the weakest contractual protection.

3.1 Subtitling, where the measurement is strongest

The SUMAT project remains the most rigorous productivity study in screen localisation. Working across seven bidirectional language pairs, which is fourteen translation directions, it rated 27,565 subtitles for quality and then ran timed post-editing productivity tasks over 37,104 subtitles with 19 professional subtitlers. Overall productivity gain was 39.90 per cent against a 25 per cent target, rising to 46.68 per cent on files where poor machine translation had been filtered out beforehand and falling to 33.12 per cent on unfiltered files. Around 57 per cent of subtitles were rated four or five on a five-point scale, meaning little or no post-editing was needed.⁹ **PEER-REVIEWED**

Two qualifications belong beside those numbers, and the second is the more important. The gains varied sharply by language pair, with English into German consistently worst on account of morphological complexity, and the authors treat automated quality metrics as an inadequate proxy for post-editing effort. And the study used statistical machine translation, predating the neural systems now in use, so the direction of the finding is robust while the specific percentages are more likely a floor than a ceiling.

WHAT HAS BEEN MEASURED, AND WHAT HAS NOT

MEASURED, AND CITABLE	MEASURED, WITH CAVEATS	NOTHING CITABLE EXISTS
Subtitling and post-editing Nine language pairs, 27,565 subtitles rated and 37,104 post-edited by 19 professionals	Metadata and shot tagging Four real projects, but no time and motion study and savings resting on testimonials	Visual effects and rotoscoping No study with a baseline and a comparison
Automated speech recognition 896 speakers across 18 accent varieties, with the error distribution reported	Assisted knowledge work A preregistered experiment with 758 consultants, though not on screen tasks	Film restoration No study with a baseline and a comparison
Synthetic voice, short utterances Two controlled listening studies with signal detection analysis	Audience response to labels Large preregistered samples, on news headlines rather than on film	Previsualisation and virtual production No study with a baseline and a comparison
		Reception of post-edited subtitles The productivity half is measured, the viewer half is not

Investment concentrates in the third column

The crafts with the strongest measured evidence are the ones attracting the least investment and carrying the weakest contractual protection. The crafts absorbing the largest share of spending have no published evidence base at all.

A craft appears in the third column where a search for baseline and comparison studies returned nothing to a defensible standard.

FIGURE 3 Screen workflows against the strength of the evidence available for them. The column on the right is where the sector is spending and where nothing citable exists.

Ecaveo assessment of the literature reviewed for this paper, references 1 and 9 to 13. **ECAVEO**

3.2 Transcription, logging and metadata

Automated metadata tagging has been tested against a human baseline, and the result is more interesting than a simple accuracy figure. In a study covering four real film projects, automated shot-type classification reached 79.5 per cent exact agreement on stills and 95.5 per cent agreement within one category, and 91.1 per cent of the dialogue-bearing audio files in one project were correctly matched to a scene. The figure that matters is the comparator. Five independent filmmakers agreed with one another only 33.5 to 65 per cent of the time on the same task.¹⁰

PEER-REVIEWED The machine was more consistent than the professionals, because the task is partly subjective and consistency is not the same as correctness.

The authors state that no direct time and motion study comparing human metadata taggers with the system is reported, and that the time savings rest on director testimonials.

3.3 Automated transcription is not uniform across speakers

Accuracy figures for automated speech recognition conceal a distribution. Testing a leading model across 18 accent varieties, researchers found match error rates of 0.115 for female speakers against 0.140 for male speakers, and 0.161 for speakers of tone languages against 0.120 for speakers of stress-accent languages, with conversational speech significantly worse than read speech.¹³

PEER-REVIEWED For a production this is an accessibility and equity question because the people whose

contributions are transcribed least accurately are systematically identifiable in advance.

3.4 Voice, where the finding is narrower than the headline

Two 2025 studies tested synthetic voice against human speech. In the larger, listeners labelled voice clones as human 58 per cent and 70 per cent of the time across experiments, against 62 per cent and 72 per cent for real human voices, with no significant difference between voice types and effectively no listener sensitivity in signal detection terms. Perceived realness was 56.6 for clones against 57.6 for humans. Clones were rated significantly more dominant.¹¹ **PEER-REVIEWED** A second study found that one commercial cloning system, and only one of the three tested, achieved naturalness and speaker similarity indistinguishable from natural human speech.¹² **PEER-REVIEWED**

Read the stimuli before reading the conclusion. Both studies used short single-sentence samples from unfamiliar speakers, pre-cleaned of breaths and clicks, with one study noting that the same clone may not convince listeners who know the voice well. Nobody has tested sustained dramatic performance, emotional range or ensemble scene work. The defensible claim is that for short, clean, isolated utterances by an unfamiliar speaker, listeners cannot tell. That is a long way from the claim that synthetic voice matches human performance, and the distance between them is where the sector's contractual argument actually sits. **PEER-REVIEWED** / **PRACTICE**

CHAPTER FOUR

Where the evidence is absent

For visual effects, rotoscoping, restoration, previsualisation and virtual production there is no peer-reviewed study with a baseline and a comparison. Not a weak one. None.

The crafts absorbing the largest share of the sector's artificial intelligence spending have no published evidence base at all. This is a strong claim and it was tested carefully before being made. A search specifically for baseline and comparison studies in visual effects, rotoscoping, film restoration, previsualisation and virtual production returned nothing to a defensible standard. What it returned instead was trade press, consultancy material and vendor blogs, none of it carrying a stated method or a baseline, together with several papers in pay-to-publish venues without credible peer review. This took longer than any other part of the paper, partly because an absence is harder to establish than a finding and you can never be sure you have looked in the right place, and partly because every promising title turned out on inspection to be a case study written by somebody with a tool to sell. **PRACTICE**

WHAT THIS RULES OUT

No productivity figure for visual effects, rotoscoping, restoration, previsualisation or virtual production should enter a budget, a business case or a board paper on the strength of published evidence, because there is none. Every number currently in circulation for those crafts originates with a company selling the tools that produce it. A vendor figure is where a local test starts.

This absence is not an argument against using these tools. Practitioners report real gains, and the absence of published measurement in a commercially sensitive field with short production cycles and heavy non-disclosure is exactly what one would expect. It is an argument about what a production may reasonably assert, and to whom. A supervisor telling a producer that a tool halves a roto pass on this show is describing experience. A slide telling a financier the same thing is describing evidence, and there is none.

4.1 Why the gap persists

Three structural reasons keep this evidence from appearing, and understanding them tells a production what it would need to do to generate its own. Production data is confidential and shot-level comparisons expose bidding rates. Shots are not comparable units, so a like-for-like baseline has to be designed in advance, and retrospective analysis will not produce one. And the people who could run the comparison are the people delivering the work, on schedules that leave no room for a control arm.

None of those is insurmountable inside a single facility on a single show, which is the argument Chapter 13 makes concrete. A facility that instruments one bounded task class across one production will know more about its own economics than the entire published literature currently contains. **ECAVEO**

4.2 The gap the audience research also leaves

One further absence is worth recording because it bears directly on localisation decisions. No peer-reviewed study measures whether viewers understand or enjoy post-edited machine-translated subtitles differently from human-translated ones. The productivity side of that trade has been measured carefully. The reception side has not been measured at all, and the decision is being taken on the half that has numbers attached to it. **PRACTICE**

That study would not be difficult to run. A distributor releasing the same title in two territories with comparable audiences already has the conditions for it, and comprehension, recall and enjoyment are measured routinely in audience research. The obstacle is that the party with the data is the party with the least interest in the answer, since a finding either way constrains a decision already taken on cost.

An absence of evidence is a position, not a neutral fact. Where nothing has been measured, the decision is taken on whatever has been measured nearby, which in this sector is almost always cost. A production that wants a different answer has to generate the evidence itself, and Chapter 13 sets out what that takes.

CHAPTER FIVE

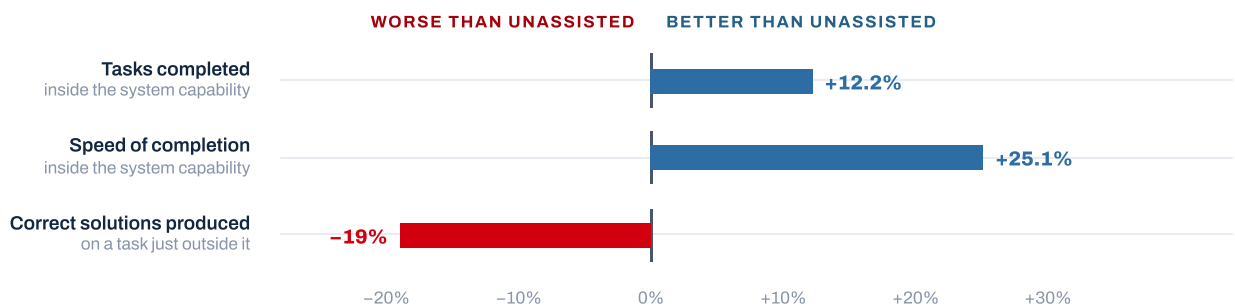
The jagged frontier

The best-designed evidence on assisted knowledge work shows gains inside a system's competence and a measurable penalty just outside it.

Because the screen-specific evidence is thin, the honest move is to reach for the best-designed general evidence and to be explicit that it is being transferred by argument. The strongest such study is a preregistered field experiment with 758 consultants, which tested performance on tasks inside and outside the capability of the system being used.

On tasks within that capability, participants completed 12.2 per cent more tasks, completed them 25.1 per cent more quickly, and delivered work of significantly improved quality. On a task selected to fall just outside it, participants using the system were 19 percentage points less likely to produce a correct solution.¹ **PEER-REVIEWED**

THE SAME TOOL, TWO SIDES OF ONE BOUNDARY



The boundary is invisible from inside the work

Work inside the capability was also of significantly improved quality. The participants could not tell which side of the boundary a task fell on, and the penalty appeared precisely where their confidence was unchanged.

FIGURE 4 A preregistered field experiment with 758 consultants. Inside the system capability, assisted work was faster, more productive and better. On a task just outside it, assisted professionals were markedly more likely to be wrong.

Reference 1. Knowledge work rather than screen production, and transferred by argument. **PEER-REVIEWED**

The finding that matters for screen production is not the size of the gain. It is that the boundary is invisible from inside the work. The consultants could not tell which side of it they were on, and the penalty appeared precisely where confidence was unchanged. A colourist, an assistant editor or a localisation manager has no better instrument for locating that boundary than the consultants had, and a production has no mechanism for detecting the crossing unless it has built one.

PEER-REVIEWED / **PRACTICE**

My own view is that this study is the strongest available argument for

bounded pilots on named task classes. A tool deployed across a department will be used on both sides of the boundary, and the organisation will see only the aggregate. The aggregate is what makes a poor result look acceptable. A tool deployed on one task class, with a baseline, produces a result that can be believed. **PRACTICE**

SPECIFICATION

What a bounded pilot has to fix before it starts

- The task class, named narrowly enough that a supervisor can say whether a given piece of work is inside it or outside it without consulting anyone.
- The baseline, measured from work already completed, because a remembered baseline flatters the tool by roughly the margin the pilot is trying to measure.
- The quality comparator, which is a blinded preference test against the unassisted version. A satisfaction score from the person who used the tool measures something else.
- The escalation route for work that turns out to sit outside the boundary, and the record that captures how often that happened, since the frequency is the finding.

The last of those is the one most often left out, and it is the one that answers the question the study actually raises. A pilot that reports only the gains has measured the easy half. The count of tasks that had to be taken back, and what they had in common, is what tells a production where the boundary runs in its own work.



PART TWO

Rights, people and audience

Five chapters on copyright that is unsettled everywhere, consent that generative tools erode by default, and a disclosure regime that damages truthful work when it is designed badly.

CHAPTER SIX

Copyright and training data, unsettled everywhere

Two judges in the same district decided two days apart in opposite directions, and the United Kingdom abandoned its preferred option in March 2026.

Nothing in this chapter or the four that follow constitutes legal advice. The material supports discussion and further review by qualified advisers and business affairs teams, and statements about how any particular production stands under these instruments are inference, and are offered as such.

A production cannot currently rely on a settled answer about the provenance of the models it is using, in any major jurisdiction. That is an uncomfortable position and it is the actual position, and the consequence is that provenance diligence and contractual warranty do the work that settled law is not yet doing.

6.1 The United Kingdom position

The Data (Use and Access) Act 2025 placed procedural duties on the Secretary of State, requiring an economic impact assessment of the four policy options in the December 2024 consultation, and a report on the use of copyright works in developing artificial intelligence systems covering technical measures and standards, data access, developer disclosure, licensing and enforcement.¹⁸ **LEGISLATION** Both were delivered on 18 March 2026.

The conclusion matters more than the process. The report states that a broad copyright exception with opt-out is no longer the government's preferred way forward, that it proposes to continue monitoring the effects of transparency rules in other countries before deciding its own approach, and that in the absence of new legislation the status quo will continue.¹⁷ **REGULATOR** There is no new text and data mining exception, no statutory transparency requirement for training data, and no announced legislative timetable. The narrow existing exception for non-commercial research is unchanged.

6.2 What the courts have and have not decided

In November 2025 the High Court of England and Wales delivered the first United Kingdom judgment on copyright works in artificial intelligence model development. The primary copyright claims were withdrawn before conclusion because the claimant could not evidence that acts of unauthorised copying had been carried out in the United Kingdom, the secondary infringement claim failed on the basis that the model itself did not constitute an infringing copy, and the trade mark

READ THIS ONE CAREFULLY

The opt-out model that dominated sector commentary through 2025 was abandoned in March 2026. Any strategy paper still describing it as the direction of travel is out of date.

claim succeeded on watermarks appearing in outputs.²² **LEGISLATION** The distinction is routinely misreported and matters a great deal. The court did not decide that training on copyright works is lawful in the United Kingdom. It decided a jurisdictional evidence question and a narrow point about what an infringing copy is.

The United States position is contradictory on its face. In June 2025 two judges in the same district reached opposite conclusions two days apart. One held that model training was transformative, described it in those terms emphatically, and held that acquiring books from pirated shadow libraries was not fair use despite the transformative end use.²⁰

LEGISLATION The other held that training on copyrighted books including shadow-library sources was fair use, found no demonstrated market substitution, declined to treat the pirated acquisition as a separate wrong, and warned that a case with stronger market-harm evidence could come out differently.²¹ **LEGISLATION** A third decision earlier that year found fair use failed as a matter of law where the use competed directly with the rights holder's own product, and is under appeal.²³

LEGISLATION

6.3 What follows for a production

My own view, which goes beyond what any of these sources states, is that the unsettled position is an argument for diligence. Waiting is not available on the timescales a production works to, and in any case the facts that will decide the question are being created now, on shows in post, by people who have never read a judgment. Three things are available to a production now and none of them requires the law to settle first. **PRACTICE**

SPECIFICATION

What to establish before a generative tool touches a deliverable

- What the supplier warrants about its training data, in the contract itself, and what indemnity sits behind that warranty.
- Whether the licence covers commercial exploitation in every territory and window the work will reach, and whether it survives a change of ownership of the supplier.
- What reference material a production is itself putting into the tool, on what rights basis, and whether that material may be retained or used for training.
- What similarity review applies to outputs before they enter a deliverable, and who performs it.

CHAPTER SEVEN

Performers, likeness and the scope of consent

A performance consent has purpose, term, territory, medium and transferability. Generative reuse erodes all five by default, and the United Kingdom has no bargained standard.

A consent to be filmed is not a consent to be reproduced, and the distance between those two things is where most of the sector's reputational risk now sits. A performance consent is bounded in at least five dimensions, and a generative system asked to produce new material from a scan or a recording is, by default, bounded by none of them.

FIVE DIMENSIONS OF A PERFORMANCE CONSENT

DIMENSION	DEFAULT IN TRADITIONAL PRODUCTION	DEFAULT ONCE A SCAN OR MODEL EXISTS	WHAT HAS TO BE WRITTEN DOWN
Purpose	The performance is used in the production it was engaged for.	A scan or a voice model can generate material for any purpose the file supports.	The named productions, cuts and marketing uses the consent covers.
Term	The engagement ends, and with it the work.	A model persists, and can be run after the engagement and after the performer's death.	An end date, and what happens to the underlying asset when it passes.
Territory	Release is bounded by the distribution agreement.	Generation is not bounded by anything, and territorial rules differ.	Which territories the consent reaches, and which it does not.
Medium	Film, television or the medium specified.	The same model serves trailers, advertising, games, dubs and social cut-downs.	The media in scope, stated one by one rather than by general words.
Transferability	Rights move with the production on ordinary terms.	A model is an asset that can be sold, licensed or acquired with the company.	Whether the consent survives a change of ownership, and on what terms.

The middle column is a default rather than an entitlement. It describes what a consent permits when nothing in it says otherwise, which is the position most United Kingdom engagements currently occupy.

FIGURE 5 Each dimension has a default in traditional production and a different default once a scan or a voice model exists. The right-hand column is what has to be written down for the consent to mean anything at delivery.

Ecaveo framework, read against references 27 and 29. **ECAVEO**

7.1 The standard that now exists, and where

The 2026 SAG-AFTRA television and theatrical agreements, ratified in June 2026 by 91.42 per cent of those voting, set the clearest bargained floor available anywhere, and seven of their provisions bear directly on the five dimensions above.²⁷ **AGREEMENT**

- Employers need an articulable business reason for scanning a performer.
- Digital alterations cannot exceed the original scripted work.
- Foreign-language dubbing by digital replica requires consent.

- Member biometric data receives specific protection.
- Replicas of minors cannot be used for nude or simulated sexual content, including by ageing up or de-ageing.
- Replicas cannot substitute for struck work.
- Protections survive a transfer of ownership.

On synthetic performers, meaning digital humans not recognisable as any specific individual, the agreement goes further than consent. Producers must bargain with the union before use, must show that the synthetic delivers significant additional value compared with what a human performer or a digital replica could achieve, and face arbitration for damages potentially exceeding what a human performer would have been paid, under a stated principle strongly favouring human performances.²⁷ **AGREEMENT** The 2023 writers' agreement established a separate and simpler principle, that no material produced by artificial intelligence can be considered literary material, that a company cannot require a writer to use it, and that the company must disclose artificial intelligence content in material given to a writer.²⁸ **AGREEMENT**

7.2 The United Kingdom has no equivalent

This is the correction with the most consequence in the paper. As at 11 September 2026 there is no concluded collective agreement on artificial intelligence between Pact and Equity, none between Pact and Bectu, and the Writers' Guild of Great Britain position is a policy statement, which carries no bargained force. The performers' union ran an indicative ballot whose results were announced in December 2025, in which 99.6 per cent of those voting supported refusing digital scanning on set, on a turnout of 75.1 per cent of 7,732 eligible film and television performers.³⁰

AGREEMENT In June 2026 the union's general secretary described the American settlement as a clear global minimum standard covering scanning, consent, data security for digital replicas, remuneration and a commitment to favour human performances, and said United Kingdom members had shown their determination to secure no less.²⁹ **AGREEMENT**

The practical consequence for a United Kingdom production is that there is no industry standard to rely on. A production writing its own consent terms is not filling a gap in the paperwork. It is taking a position in a live industrial dispute, and it should know that it is doing so. My own view is that a producer who writes terms at or above the American floor now will be in a materially better position than one who waits, both contractually and with talent. **PRACTICE**

ILLUSTRATIVE SCENARIO**The scan that outlived the schedule**

An invented drama scans its principal cast during pre-production for a sequence involving crowd duplication. The scan is taken under a standard release drafted before generative tools existed, which assigns the producer all rights in the material captured, in all media, in perpetuity.

Two years later the series is recommissioned, one performer has declined to return, and a post facility proposes generating three lines of dialogue from the existing scan and voice recordings to bridge a narrative gap. The release appears to permit it. The performer's agent, the union, the broadcaster's compliance team and the territory with a synthetic-media disclosure rule each ask a different question, and the production discovers that a document drafted for one purpose is being asked to carry four.

Nothing in this scenario requires bad faith. It requires only a consent whose scope was never specified, which is the ordinary case.

CHAPTER EIGHT

The workforce, and which crafts are exposed

The exposure forecasts are scenario models commissioned by interested parties, and the craft they place at greatest risk is the one where the automation evidence is strongest.

Two studies dominate the workforce conversation and both should be read with their commissioning and method visible, because both are frequently quoted as measurement when neither is.

The first surveyed 300 senior leaders across six entertainment industries between November and December 2023 and combined the results with occupational modelling. Three quarters of respondents indicated that generative artificial intelligence tools had supported the elimination, reduction or consolidation of jobs, and the modelling projected that 21.4 per cent of film, television and animation roles, about 118,500 jobs, would be affected by 2026. It was commissioned by an animation union, a concept art association, an artists' campaign and a cartoonists' foundation.³¹ **SURVEY** The 2026 horizon has now arrived and no follow-up validating the projection appears to exist, which is itself informative.

The second modelled revenue at risk across music and audiovisual creators to 2028, drawing on expert interviews and workshop sessions, and projected annual revenue loss for audiovisual creators of €4.5bn by 2028, representing 21 per cent of total audiovisual creator revenues. It was commissioned by an international confederation of authors' societies.³² **REVIEW** The report states that unanticipated events may occur and that some assumptions may not be realised. These are scenario forecasts built on expert elicitation, not observed losses.

8.1 The figure inside the forecast that deserves attention

Underneath the headline number sits a breakdown by craft, and it is far more useful than the total. Translators and adapters face a projected 56 per cent cannibalisation of their revenues by 2028, against a band of 15 to 20 per cent for screenwriters, directors and co-authors.³² **REVIEW**

Set that next to Chapter 3. The craft with the highest projected exposure is the same craft with the strongest published evidence that automation works, and it is the craft with the least contractual protection and the least public sympathy. My own view is that a sector serious about responsible adoption would address localisation first, precisely because it is where the argument is already lost on the evidence and where the terms are being set by default. **PRACTICE**

8.2 What the United Kingdom data does and does not show

United Kingdom employment data on artificial intelligence specifically does not exist in citable form. The largest recent workforce survey, covering 5,597 respondents with more than 3,600 from television, found around 45 per cent of television workers out of work in March, at 45 per cent of drama crew, 46 per cent of unscripted workers and 45 per cent of commercials staff, with fewer than one in five reporting employment back to pre-2023 levels.³³ **SURVEY** That survey does not address artificial intelligence and attributes the contraction to post-strike conditions. Any causal link between United Kingdom screen unemployment and artificial intelligence is, on the published evidence, unevidenced, and a responsible paper should decline to draw it. The honest position is that nobody knows, and the people best placed to find out have their own reasons for not asking the question in a form that would produce an answer. **PRACTICE**

8.3 The cost nobody is counting

One further exposure sits outside the employment debate and has no measurement mechanism at all. The screen sector's principal carbon calculator, in its 2025 methodology, contains no treatment of artificial intelligence, cloud computing or data centre emissions. It measures editing-suite energy without distinguishing equipment types and does not account for cloud-based post-production services.³⁴ **REGULATOR** Against that, the best available measurement finds image generation using on average over 60 times more energy than text generation, at 2.907 kilowatt hours against 0.047 per thousand inferences, with the least efficient image model consuming the equivalent of 522 smartphone charges.¹⁴ **PEER-REVIEWED** A production committed to a carbon standard that cannot see its own compute is reporting a number it knows to be incomplete, and the remedy is to count it locally until the standard catches up. **PRACTICE**



Virtual production concentrates decisions that were previously separated by months, which raises both the value of getting the record right and the cost of getting it wrong.

CHAPTER NINE

Disclosure, and why the flat label backfires

A bare statement that artificial intelligence was used depresses trust in truthful work. A statement of what the system did largely does not.

Disclosure is the sector's instinctive remedy and audiences plainly want it. In a nationally representative British survey of 2,035 adults in October 2025, 91 per cent said platforms should be required to label artificial intelligence generated content clearly, 53 per cent said they would not engage with content created or co-created by artificial intelligence against 26 per cent who would, and the leading concerns were authenticity at 67 per cent, misleading content at 62 per cent and job displacement at 62 per cent.⁴ **SURVEY** Only 7 per cent reported no concerns at all. An earlier report, surveying 2,000 British adults between 16 December 2024 and 2 January 2025, found 86 per cent saying it is important to disclose generative artificial intelligence use.⁶ **SURVEY**

Audiences also do not trust the labels they are asking for. In an earlier nationally representative survey of 2,128 United Kingdom adults, 48 per cent said they would distrust the accuracy of artificial intelligence content labelling against 19 per cent who would trust it, while half thought labels would be effective at reducing misinformation.⁵ **SURVEY** Wanting a label and believing it are different things, and a disclosure regime has to survive both.

9.1 The finding that should change the design

The most consequential evidence in this paper comes from two preregistered experiments with 4,976 participants. Labelling a headline as artificial intelligence generated reduced perceived accuracy and willingness to share by about 0.17 points on a six-point scale in the first experiment and 0.11 in the second, and, critically, the effect was significant and of similar size across true and false headlines. Human-made headlines mislabelled as artificial intelligence generated suffered a comparable penalty. Scepticism appeared where respondents had no definition of the system's role or a strong definition in which it selected the topic and wrote everything. A weak definition, in which the system improved clarity and style, produced minimal scepticism.¹⁵

PEER-REVIEWED

THE DISCLOSURE LADDER

A statement of function

"Dialogue was processed for linguistic accuracy, and the performance was unaltered"

An element-level statement

"This sequence contains generated elements, and these are they"

A strong definition

"The system selected the subject and produced the material"

A bare label

"Artificial intelligence was used in the making of this film"

MINIMAL MEASURED SCEPTICISM

A weak definition, in which the system improved clarity and style, produced very little penalty in the same experiments.

MOVES THE QUESTION FROM THE WORK TO THE ELEMENT

Untested directly. Proposed here because it supplies the definition the evidence shows is doing the work.

SCEPTICISM COMPARABLE TO THE BARE LABEL

Respondents given this description responded much as those given no description at all.

LARGEST MEASURED PENALTY

Perceived accuracy and willingness to share fell by about 0.17 points on a six-point scale, and the fall was of similar size on true and false material.

Rising specificity about what the system did, and the measured audience response at each level

FIGURE 6 What the disclosure says the system did drives audience response, and the bare fact of disclosure does not. A bare label carries a penalty that falls on truthful work as heavily as on false work.

Ecaveo framework, built on references 15, 16 and 8. **ECAVEO**

A companion study of 7,579 participants found that process-based labels reading artificial intelligence generated reduced belief less than harm-based labels reading manipulated or false, and did not significantly change stated sharing or liking behaviour. Its authors raised the possibility that labelling a subset of content might boost the credibility of unlabelled content, and explicitly did not measure it.¹⁶ **PEER-REVIEWED** That question remains open and matters for any sector-wide scheme.

The design principle that follows is specific. Disclose what the system did rather than that a system was used. A production that states artificial intelligence was used in the making of this film has chosen the form of words with the worst measured effect on audience trust in its own truthful work. A production that states which element was generated, or that dialogue was processed for linguistic accuracy without altering the performance, has chosen the form with the least. **PEER-REVIEWED** / **ECAVEO**

9.2 Objection is to specific uses, not to artificial intelligence

French public research found acceptance varying sharply by use, with entirely generated works and virtual actors largely rejected, dubbing and screenwriting mixed, and use in visual effects widely accepted, alongside a majority wanting to be informed.⁸ **SURVEY** A comparable gradient appears in news research across six countries, where comfort rose from 12 per cent for content made entirely by a machine to 43 per cent where a human leads with some assistance and 62 per cent for entirely human work.⁷ **SURVEY** These are news audiences. The gradient is likely to transfer, and the absolute numbers should not be quoted as screen figures.

9.3 What the rules actually require

Article 50 of the European Union artificial intelligence regulation applies from 2 August 2026. Providers of systems generating synthetic audio, image, video or text must ensure outputs are marked in a machine-readable format and detectable as artificially generated or manipulated, and deployers of systems producing deep fakes must disclose that the content has been artificially generated or manipulated.¹⁹ **LEGISLATION** For film the operative provision is the second subparagraph of Article 50(4), which provides that where content forms part of an evidently artistic, creative, satirical, fictional or analogous work, the transparency obligations are limited to disclosure of the existence of such generated or manipulated content in an appropriate manner that does not hamper the display or enjoyment of the work.¹⁹ **LEGISLATION**

Two points of timing are commonly misstated. The Digital Omnibus regulation that entered into force in July 2026 postponed high-risk obligations to December 2027 and August 2028, and did not delay Article 50. It introduced one transitional measure, giving systems placed on the market before 2 August 2026 until 2 December 2026 to comply with the marking obligation, with no grace period for anything placed on the market afterwards.²⁴ **LEGISLATION** A voluntary Code of Practice on transparency of artificial intelligence generated content was published in final form in June 2026, with accompanying guidelines in July 2026.²⁵

REGULATOR

In the United Kingdom there is no equivalent. The broadcasting regulator's 2026 strategic approach proposes monitoring and a shared evidence base, and stops short of a code amendment. Its regulation is framed as technology-neutral and concerned with outcomes for audiences, whichever technology produced them.³⁹ **REGULATOR** In advertising, the position is that there is no blanket legal requirement in the United Kingdom to disclose the use of artificial intelligence in ads, that the test is whether audiences would be misled without disclosure, and that disclosure alone is very unlikely to mitigate the harm caused by a fundamentally misleading message.⁴⁰ **REGULATOR**

CHAPTER TEN

Provenance, and what marking cannot do

In the regulator's own testing, cropping an image removed the watermark. Provenance is a discipline of record, not a property of a file.

The sector's technical answer to disclosure is provenance marking, and it is a good answer to a narrower question than the one usually asked of it. The content provenance standard now at version 2.4, published in April 2026, attaches signed assertions describing an asset's origin and edit history, and is backed by a steering committee including broadcasters, platforms, camera manufacturers and model developers.³⁵ **STANDARD** Used inside a controlled pipeline it does real work. There is a version of this argument that treats marking as a solved problem waiting only for adoption, and I held something close to it until I read the regulator's own results.

It does not survive contact with distribution, and that failure is documented. An independent technical review identifies metadata stripping through file uploads as the primary obstacle to interoperability, notes that the framework measures the trustworthiness of the signer and says nothing about the data, records that absence of a credential can itself count against content, and documents cases where fraudulent material carried valid credentials.³⁶ **REVIEW**

WHAT SURVIVES, AND WHAT DOES NOT

THE MARK SURVIVED

Small erasures

Localised removal of detail left the mark detectable.

Text overlays

Burned-in captions and titles left the mark detectable.

Colour alterations

Grading adjustments left the mark detectable.

THE MARK DID NOT

Cropping

The regulator records that simple adjustments such as cropping an image resulted in the removal of the watermarks.

Platform upload

Metadata stripping through file uploads is identified as the primary obstacle to interoperability.

Regeneration attack

Adding noise and reconstructing the image beat four pixel-level schemes on both detection rate and image quality.

Method. Three well-known and openly available watermarking tools, hundreds of watermarked images, 26 distinct edit types applied across approximately 100 variations at differing intensities, with detection tested after each. The first two rows on the right are incidental losses rather than attacks, which is what makes them the harder problem.

FIGURE 7 Results of the regulator's own watermarking experiment, which applied 26 distinct edit types across roughly 100 variations to hundreds of watermarked images using three openly available tools. Deliberate removal is a separate and easier problem.

Reference 37. **REGULATOR**

The regulator ran its own experiment. Taking three well-known and openly available watermarking tools, it watermarked hundreds of images, applied 26 distinct edit types across approximately 100 variations at differing intensities, and tested whether the watermark remained detectable. Small erasures, text overlays and colour alterations were survivable. Its finding on the simplest edit of all is the one worth carrying, namely that simple adjustments such as cropping an image resulted in the removal of the watermarks.³⁷ **REGULATOR**

Deliberate removal is easier still. A peer-reviewed regeneration attack, which adds noise to destroy the watermark and then reconstructs the image, consistently achieved lower detection rates and higher image quality than previous techniques against four pixel-level schemes, and its authors argue for a shift from invisible watermarks towards semantic-preserving ones.³⁸ **PEER-REVIEWED**

WHAT THIS RULES OUT

No reliance on watermarking or embedded credentials as an enforcement mechanism, a compliance control or a substitute for the production's own record. No treatment of a missing credential as evidence that an asset is synthetic, or of a present one as evidence that it is authentic. And no disclosure strategy whose operation depends on metadata surviving a platform upload. The published evidence is that it frequently does not.

My own view is that this reframes provenance. It does not defeat it. Marking is a convenience that helps assets carry their history within a pipeline the production controls. The obligation it is often asked to discharge, which is proving at delivery what an asset is made of, has to be met by a record the production holds independently of the file. That record is the subject of Chapter 11. **ECAVEO**



PART THREE

Building the capability

Six chapters on what to record, what to tier, what to pilot and what to refuse, and on the manifest that makes each of those decisions answerable at delivery.

CHAPTER ELEVEN

The asset manifest

Seven fields, captured when the asset is made, that answer every question asked about it at delivery.

The manifest is the smallest record that makes the clearance cliff survivable. It is deliberately short, because a record that takes an artist more than a minute to complete will not be completed, and a record that is not completed is worse than none, since it creates the appearance of a control that does not operate. **ECAVEO**

THE ASSET MANIFEST

	FIELD, CAPTURED WHEN THE ASSET IS MADE	THE DELIVERY QUESTION IT ANSWERS
1	Asset identifier, and the shot or cue it belongs to	Which deliverables and versions is this in?
2	Source material, and where each element came from	What is this frame made of?
3	Rights basis for that material, rather than a rights conclusion	On what basis were we entitled to use it?
4	Consent reference, where a person is represented	Does the consent reach this use, medium and territory?
5	Tool and version, and the account it ran under	Can the supplier warranty be matched to the asset?
6	Transformation applied, in one line of plain description	What did the system actually do?
7	Named reviewer, and the date of review	Who checked it, and can they be asked?

A record rather than a rule

Every field is information somebody already holds at the moment the asset is made. The difficulty is that they hold it in their head, in a message thread, or in a system that will be decommissioned when the show wraps.

FIGURE 8 Seven fields, captured at creation and carried with the asset through vendors, versions and windows. The left column is what is recorded; the right is the delivery question it answers.

Ecaveo framework. **ECAVEO**

Three design decisions inside it are worth stating. Organisations tend to get all three wrong. The manifest records the rights basis rather than a rights conclusion, since a conclusion reached in 2026 may not survive 2028 while the underlying facts will. It records the consent reference, so that a performer's agreement is held once and cited many times, and never copied into systems it should never reach. And it names the reviewer, because a department cannot be asked what it was looking at.

11.1 Where it is captured

The manifest is worthless as a separate database that somebody is supposed to update. It has to be attached to whatever the production

already treats as the spine of its asset tracking, which in most cases is the editorial or production asset management system, and it has to be a required field on the handoff between vendors. A facility that cannot deliver a manifest with a shot should be told that at the bidding stage and not at delivery.

Nothing in it requires new technology. Every field is information that somebody in the production already holds at the moment the asset is made, and the entire difficulty is that they hold it in their head, in a message thread or in a system that will be decommissioned when the show wraps.

11.2 What it is not

A manifest is evidence, and evidence is all it is. It is what a disclosure decision gets taken from, by the people entitled to take it, at the point at which the question actually arises. Chapter 10 explains why it is no substitute for a provenance credential, and it is not itself a disclosure. A production that holds manifests can answer a platform, a regulator, an insurer, a union or a journalist. A production that does not is reconstructing from memory, under time pressure, and usually in public.

CHAPTER TWELVE

An operating model

Ordinary controls, tiered by what has to be true at delivery rather than by which tool was used.

The controls in this chapter are ordinary. Their value lies in being specified before a tool is used, and in being tiered by consequence. Assembled after a question has been asked, the same controls are an explanation.

12.1 A tier for every use

The tiers below are ordered by clearance exposure, meaning what has to be true at delivery, in every territory and window the work will reach, for the asset to be usable. Data sensitivity and performer involvement operate as modifiers that can raise a use by one tier and never lower it.

ECAVEO

RISK TIERS ORDERED BY CLEARANCE EXPOSURE

TIER	WHAT HAS TO BE TRUE AT DELIVERY	MINIMUM CONTROLS	APPROVAL
One	Nothing. The output never reaches a deliverable, a participant or a third party, and can be discarded without trace.	Approved tool, user training, source checking, no external use without review.	Department head
Two	The output informs an internal decision or reaches colleagues, and a correction reaches everyone who saw it.	Registered use, data classification, named reviewer, logging, no unreleased material in uncontracted systems.	Head of department with technology and, where relevant, legal
Three	The asset is delivered, published or licensed. Rights, consent and disclosure must be answerable in every territory and window.	Asset manifest, supplier warranty and indemnity, similarity review, disclosure decision, territory check, recorded creative sign-off.	Producer with business affairs, legal and the accountable creative owner
Four	A performer's likeness or voice is generated or altered, or an asset is presented as something other than what it is.	All tier three controls, plus bounded specific consent, union and agreement check, estate or representative approval where applicable, and a recorded editorial decision on disclosure.	Executive sponsor with legal, business affairs, workforce representation and communications

12.2 The risks that actually recur

Six failure modes account for most of what goes wrong, and each has a guardrail that is cheap in advance and expensive afterwards.

SIX FAILURE MODES, AND THE GUARDRAIL FOR EACH

RISK	FAILURE MODE	REQUIRED GUARDRAIL
Training provenance	Inputs, references or outputs reproduce protected expression, or the model rests on training data the supplier will not describe.	Approved tool and model register, contractual warranty with indemnity, licensed inputs, similarity review before an output enters a deliverable.
Likeness, voice and biometrics	A scan, performance or replica is created, reused or transferred beyond what was actually agreed.	Bounded specific consent stating purpose, term, territory, medium and transfer, with security, revocation and a named holder of the record.
Confidential material	Scripts, rushes, cuts, casting data or commercial terms enter a system whose retention and training terms are uncontracted.	Production data classification, approved tenancy, restricted connectors, contracted retention and deletion, incident duties.
Authenticity and disclosure	A synthetic or materially altered element misleads an audience, or cannot be traced when a platform or regulator asks.	Asset manifest, disclosure decided by function in pre-production, territory release check, editorial ownership of the decision.
Creative homogenisation	Outputs reproduce stereotypes, exclude communities, or collapse distinctive work toward a familiar mean.	Representative review, cultural expertise on the specific material, diverse test sets, and a recorded creative challenge that somebody is empowered to make.
Workforce and compute	Automation removes career-forming work, or compute use sits outside a carbon commitment the production has made.	Consultation before deployment, a skills route for affected crafts, a stated preference for human work, and local measurement of compute.

12.3 Technical controls that are not negotiable

- Production material is classified before any tool touches it, and unreleased scripts, rushes, cuts, casting data and commercial terms never enter a system whose retention and training terms are not contracted.
- Contracts define training use, retention, hosting region, sub-processors, deletion, security incident duties and audit rights, in the agreement itself.
- Identity and project permissions are enforced before retrieval and before tool use, because a model is not a rights clearance and cannot be made into one.
- Agents receive least-privilege, short-lived credentials, consequential actions require explicit approval, and every action is logged and reversible.
- Every production system records model, version, prompt, retrieval corpus, tools, permissions and evaluation version, so that a delivery question can be answered and an incident investigated.
- Prompt injection, malicious documents, sensitive information disclosure and insecure output handling are tested before scale and after material change.
- Generated work passes the same review, approval and records controls as work produced without assistance.

These are consistent with the national cyber security guidance on secure artificial intelligence system development, the generative artificial

intelligence risk profile published by the United States standards institute and the information regulator's accountability expectations.^{41,42,43} **REGULATOR** None is novel. The discipline lies in applying them to the fourth vendor on the fourth show as rigorously as to the first.

12.4 The register, and who keeps it

One artefact holds the tiers together. It is the least glamorous thing in this paper. A production needs a single list of the uses it has approved, the tier each sits in, the person accountable for each and the date each was last looked at. Without it the tiers are only a taxonomy, because nobody can say which tier a given use was placed in or who decided.

The register belongs to the production, for the same reason the manifest does. A studio-level register records what a company permits in principle. A production-level register records what this show is actually doing, which is the question a platform, an insurer or a union will ask. Where both exist the production register governs, and the studio register is the standard it has to meet. **ECAVEO**

Two habits keep it honest. Every entry carries a review date, so that an entry which has not been looked at in a year is visible as such. And a use that has been refused is recorded alongside the ones approved, with the reason, because the refusals are what tell a new vendor where the line sits without another round of argument. Whether a register survives into its second year is the real test of whether any of this was ever more than paperwork, and I have not seen enough of them running that long to say.

THE SHORTEST USEFUL VERSION

One row per approved use, with tier, owner and review date. A register that nobody can read in five minutes will not be read at all.

CHAPTER THIRTEEN

Evaluation and assurance

A facility that instruments one bounded task class on one production will know more about its own economics than the published literature contains.

Chapter 4 established that no published evidence exists for the crafts where the sector is spending most. The corollary is that a production or facility willing to measure one thing properly acquires an advantage that cannot currently be bought, because nobody else has the number either.

ECAVEO

Each pilot should open with a one-page protocol written before anyone is given access, recording the current workflow, the target outcome, the sample, the comparison method, the quality rubric, the plausible harms, the review burden and the conditions under which the pilot stops. Adoption figures and user enthusiasm are evidence of interest rather than evidence of value, and a pilot that reports only usage has answered a question nobody needed answering.

SPECIFICATION

Five design choices that cover most screen workflows

- Paired shots or paired reels, where comparable work can be completed with and without assistance without contaminating the comparison.
- A staggered rollout across units or episodes, where withholding a tool indefinitely is impractical.
- Blinded review by a supervisor or colourist who does not know which version was assisted.
- Adversarial testing covering prompt injection, permission boundaries, confidential material and unsafe tool use, before the tool reaches unreleased content.
- Repetition after any material change to the model, the prompt, the corpus or the workflow, because a tool evaluated once was evaluated against a system that no longer exists.

13.1 What to measure, and what would stop a rollout

WHAT TO MEASURE, AND THE GATE EACH MEASURE CONTROLS

DIMENSION	EXAMPLE MEASURES	GATE BEFORE SCALING
Outcome	Turnaround, iterations to approval, search and retrieval time, shots or minutes delivered.	A predefined improvement over a captured baseline, reported with a range rather than a single figure.
Quality	Defect and correction rate, reviewer notes, rework passes, blinded preference against the unassisted version.	No material quality regression, and a critical-error rate below the agreed threshold.
Craft judgement	Override behaviour, escalation quality, whether the operator can explain what the tool did and where it fails.	Practitioners retain independent judgement and can locate the boundary of the tool.
Rights and clearance	Manifest completeness at delivery, unanswerable clearance questions, similarity review findings.	Every delivered asset answerable without contacting the artist who made it.
Security and confidentiality	Permission leakage, unreleased material exposure, unsafe actions, red-team findings.	No unresolved critical finding, and every high finding mitigated or formally accepted by a named owner.
Economics	Licence, compute, integration, assurance, review and support cost per delivered unit.	A sustainable total cost set against the verified benefit rather than the projected one.
Operations	Reliability, vendor change, monitoring, support route and rollback.	A named service owner and tested operational controls in place before scale.

13.2 The minimum assurance record

A use that cannot produce the following record has not been governed, whatever policy it complied with. It sits above the asset manifest, which records individual assets, and describes the system that produced them.

THE MINIMUM ASSURANCE RECORD FOR A DEPLOYED SYSTEM

RECORD FIELD	REQUIRED CONTENT
Purpose and owner	The problem, the intended users, the benefit claimed, the accountable owner and the review date.
Boundaries	Permitted and prohibited uses, data classes, territories, asset types and actions.
System description	Model and version, vendor, prompts, retrieval sources, tools, interfaces and dependencies.
Contractual position	Training use, retention, hosting region, sub-processors, deletion, incident duties, audit rights, warranty and indemnity.
Evaluation	The sample, the method, the results with their limitations, and the approval decision.
Operations	Monitoring, incidents, user feedback, change control, rollback and retirement.

CHAPTER FOURTEEN

A portfolio built to the evidence

Localisation moves up because it is measured. Generative video moves down because it is not, and because the exposure lands at delivery.

The portfolio below revises the conventional opportunity list in the light of Parts One and Two, and each revision has a reason attached to it that comes out of the evidence. Archive, metadata and localisation move up, because they have measured evidence and contained exposure. Generative elements in delivered picture move down, because the evidence is absent and the bill arrives at delivery. Anything touching a performer's likeness carries a consent and agreement gate, and what the technology can do is beside the point. **ECAVEO**

A CANDIDATE PORTFOLIO, ORDERED BY THE STRENGTH OF THE EVIDENCE BEHIND IT

USE	VALUE HYPOTHESIS	TIER	POSITION
Transcription, logging and metadata tagging	Searchable material and recovered assistant time.	One to two	Start here. Measured against a human baseline, more consistent than the professionals it assists, and contained within cleared material.
Archive discovery and catalogue enrichment	Higher catalogue value and reusable assets found rather than reshot.	Two	Strong fit. Retrieval over material the organisation already owns, with rights status carried in the index rather than inferred.
Subtitling by post-edited machine translation	Broader reach at lower cost per language.	Two to three	The best-evidenced productivity case in the sector, and the one where terms should be negotiated rather than defaulted. Reception evidence is absent.
Development research over licensed sources	Better-informed development with authorship unchanged.	Two	Citation-required retrieval over permissioned archives only. The tool prepares the evidence and the executive selects.
Production coordination and breakdown assistance	Less scheduling friction and faster iteration.	Two	Department-head validation on every output, and no autonomous decision on safety, employment or spend.
Post-production clean-up and restoration	Shorter cycles on bounded, repetitive passes.	Three	Pilot with a baseline, because no published evidence exists. Shot-level manifest from the first frame, and a defect measure alongside turnaround.
Localisation dubbing using synthetic voice	Faster versioning and wider access.	Three to four	Only with performer consent bounded by purpose, term, territory and medium, and with the disclosure position decided before delivery.
Generative elements in delivered picture	New imagery at lower cost.	Three to four	Defer beyond bounded, recorded uses. No published evidence, unsettled training-data law, and the exposure surfaces at clearance rather than in the suite.
Digital replica of a performer	Continuity, de-ageing, or completion of unfiled material.	Four	Outside an adoption programme. Route to legal, business affairs, the union position and, for a deceased performer, the estate, from the outset.
Synthetic performer	A digital human not recognisable as any individual.	Four	Defer. The bargained position elsewhere requires prior negotiation and a significant additional value test, and no United Kingdom equivalent exists.
Autonomous agents with write access to production systems	Fewer manual coordination steps.	Four	Defer broad autonomy. Permit bounded actions with least privilege, approvals, logging and tested rollback.

14.1 Three pilots worth starting with

The first should attach a manifest to every artificial intelligence touched asset on one production. Begin on the first day of post. It puts the record and the vendor handoff under load, and it asks the delivery question a year early, without requiring any new tool at all. It also produces the only artefact in this paper that has value on every subsequent show. Success is measured by manifest completeness at delivery and by how many clearance questions can be answered without contacting the artist who did the work.

The second should test one bounded post-production task class on cleared material, with a baseline, blinded review and a defect measure. Transcription and logging, metadata tagging or a defined clean-up class are the natural candidates because the comparator is knowable. Given the finding in Chapter 5, the protocol must include a class of work the tool is expected to handle badly, because a pilot that tests only the favourable case measures the vendor's claim and tells the facility nothing about its own economics.

The third should build the localisation evidence the sector is deciding without. Post-edited machine translation has measured productivity gains and no measured reception evidence, so a production comparing comprehension and viewer response between post-edited and human-translated subtitle tracks on the same title would close a genuine gap. This is the least glamorous of the three and the one most likely to change a commissioning decision.



Localisation and access is where the evidence is strongest, the exposure highest and the contractual protection weakest. A sector serious about responsible adoption would start here.

CHAPTER FIFTEEN

Delivery over six to nine months

Four phases, three pilots and one manifest, arranged so that the organisation owns the record when the engagement ends.

A PHASED PROGRAMME OVER SIX TO NINE MONTHS

PHASE	TIMING	PRIMARY OUTPUTS	EXIT CONDITION
Discover	Weeks 1 to 4	Workflow observation across development, production, post, localisation and archive; tool and undeclared-use inventory; rights and consent audit; territory map; baseline measures.	The executive agrees priorities, risk appetite and pilot owners, and one delivered asset has been traced backwards to establish what can be proved.
Prove	Weeks 5 to 10	Manifest on one production; one bounded post task class with a baseline and blinded review; one localisation comparison; supplier and contract review; first adversarial tests.	At least two pilots show value without a material quality or rights regression, and the manifest survives a vendor handoff.
Standardise	Weeks 11 to 18	Use intake, clearance tiers, consent scope templates, manifest specification in the asset management system, supplier clauses, disclosure rules and role training.	A new production can start low-risk work through a clear, repeatable route without the programme team.
Scale	Months 5 to 9	Expanded patterns across productions and catalogue, show-level owners, portfolio assurance, benefits tracking and a compute measure.	Scaled uses meet their gates, named owners accept ongoing operation, and the record survives the end of the engagement.

15.1 The first 30 days

Observe real work across development, production, post, localisation and archive before proposing anything, and interview business affairs, legal and the head of post in the first week, because they hold the delivery questions the rest of the programme has to answer. Inventory the artificial intelligence already present, including features embedded in software the organisation already licenses, since most organisations discover at this point that adoption began without them. Map the data boundaries, the contractual duties, the territories the work will reach and the named people who actually decide. A word on the shape of the thing. Programmes of this kind tend to fail in about month four, when the person carrying the work on top of a full job runs out of evenings, so the phasing below is built around what one person can actually hold.

PRACTICE Publish interim safe-use guidance and a simple route to propose an experiment without blame, on the reasoning that undeclared use is a governance failure and punishing disclosure guarantees more of it. Select three pilots with baselines, representative test material and stop conditions. And run one bounded workflow end-to-end in a sandbox, so that leaders see the capability and the control design in the same sitting.

15.2 Operating rhythm

A weekly delivery review tracks pilot evidence, risks and blockers. A fortnightly clinic with legal, business affairs and technology resolves intake questions quickly, since the main cause of undeclared use is a governance process slower than the work it governs. A monthly executive review makes scale, redesign or stop decisions using one evidence template. Anything touching a performer, a participant or a released asset adds the relevant creative, legal, workforce and communications voices before work begins.

THE MONTHLY DECISION, AND THE EVIDENCE EACH OPTION REQUIRES

DECISION	EVIDENCE NEEDED
Scale	The target benefit achieved against a captured baseline, quality maintained under blinded review, residual rights and clearance risk formally accepted, and ownership funded.
Redesign	A value signal exists, and quality, review burden, workflow fit, manifest completeness or controls fall short of the gate.
Stop	No verified benefit, unacceptable clearance or rights exposure, poor fit with how the craft actually works, vendor dependency, or a simpler method performing better.

15.3 What success looks like at the end

WHAT SUCCESS LOOKS LIKE, AND HOW IT WOULD BE EVIDENCED

OUTCOME	EVIDENCE
A provable catalogue	Delivered assets whose composition can be established without contacting the person who made them.
Verified value	A small portfolio of scaled workflows with measured improvement against a baseline captured beforehand.
Consent that means something	Performer terms drafted for the tools that exist, bounded in purpose, term, territory, medium and transfer.
A disclosure position	Decided in pre-production by the people who own the creative and legal risk, and stated by function rather than by fact.
Durable ownership	Business affairs, post and technology maintain the manifest and the controls after the programme ends.
A credible next step	A backlog that separates proven scale opportunities from research, and both from what the organisation has decided to refuse.

CHAPTER SIXTEEN

The intelligent production programme

What to adopt, in what order, and what to measure. The first step costs nothing and is the one everybody skips.

The adoption sequence that follows from this paper is deliberately unimpressive, which is the point. The sector has spent three years acquiring capability and has not yet built the record that would let it prove what that capability produced.

16.1 The sequence

First, attach a manifest to one asset class on one production. Nothing else in this list works without it. Second, write consent scope before anyone is scanned, covering purpose, term, territory, medium and transfer. Third, tier by clearance exposure. Fourth, decide disclosure by function, and decide it in pre-production. Fifth, baseline one task class properly, including the cases the tool should handle badly. Sixth, and only sixth, consider anything generative in delivered picture, and expect to conclude that the first five have already produced most of the value.

ECAVEO

16.2 Questions a board should ask before approving any of this

- Which of our current productions could produce, today, a record of what a given delivered asset is made of?
- Which artificial intelligence is already present in software we license, and who approved it?
- What do our performer releases actually permit, and were they drafted before or after generative tools existed?
- Which territories and platforms will ask us a synthetic media question on this title, and who is preparing the answer?
- Where does our vendor chain break the chain of custody, and what does the contract say about it?
- What have we measured ourselves, and what have we accepted from a supplier?
- What will we decline to do even when it becomes technically straightforward and commercially attractive, and who has the standing to hold that line?
- What evidence would tell the board that adoption is increasing the value of the catalogue, and what would tell it the risk is growing faster?

16.3 What good looks like in three years

A production that can say what a frame is made of without ringing the

artist who made it. Consent terms written for the tools that exist now. A disclosure position decided in pre-production by the people who own the creative and legal risk. And a portfolio in which the uses that reach an audience are the best evidenced ones, which is the reverse of where the sector starts.

That is a modest description of progress in a field that has been sold something far more exciting for three years. It is also, on the evidence assembled here, most of what is actually available.

The technology is not the thing that arrives late. The question is. By the time anyone asks what a frame is made of, the only thing that helps is a record somebody thought to keep.

CHAPTER ONE

END MATTER

Board and executive checklist

Eighteen items. The first four cost nothing and answer most of what is asked at delivery.

- | | |
|---|--|
| ☐ Every AI-touched asset carries a manifest
Sources, rights basis, consent reference, tool and version, transformation, reviewer and permitted uses. | ☐ The manifest is a required handoff field
Not an optional one, and raised at bidding. |
| ☐ Consent scope is written before anyone is scanned
Purpose, term, territory, medium and transferability, each one stated on the face of the document. | ☐ Disclosure is decided in pre-production
By the people who own the creative and legal risk, not during a campaign. |
| ☐ Disclosure states function, not fact
What the system did, because a bare label depresses trust in truthful work. | ☐ Tiering is by clearance exposure
What has to be true at delivery, in every territory and window the work will reach. |
| ☐ Unreleased material never enters an uncontracted system
Scripts, rushes, cuts, casting data and commercial terms. | ☐ Supplier training-data warranties are in the contract
With an indemnity behind them, not in the marketing material. |
| ☐ Licences cover every territory and window
And survive a change of ownership of the supplier. | ☐ Similarity review runs before an output enters a deliverable
By a named person, against a defined standard. |
| ☐ No model is treated as a rights clearance
Permissions are enforced before retrieval, not inferred from output. | ☐ Watermarking is not relied on as a control
The regulator's own testing shows a crop removes it. |
| ☐ A missing credential proves nothing
And a present one proves less than it appears to. | ☐ One task class is baselined properly
Including the cases the tool is expected to handle badly. |
| ☐ Blinded review is used where the judgement is subjective
Because the reviewer's expectation is part of the measurement. | ☐ Localisation terms are negotiated, not defaulted
It is the craft with the strongest automation evidence and the weakest protection. |
| ☐ Compute is counted
The sector's carbon methodology does not yet capture it, so the production must. | ☐ Evaluations repeat after material change
A tool evaluated once was evaluated against a system that no longer exists. |

END MATTER

Glossary

Terms used in a specific sense in this paper.

Clearance cliff

The gap between the moment an AI decision is taken and the moment it is examined at delivery, clearance, insurance or release. The organising problem of this paper.

Consent scope

The five bounded dimensions of a performance consent, namely purpose, term, territory, medium and transferability.

Synthetic performer

A digital human not recognisable as any specific individual, subject to a separate bargaining obligation under those agreements.

Content credential

A signed, machine-readable assertion about an asset's origin and edit history under the provenance standard.

Jagged frontier

The invisible boundary of a system's competence, either side of which assisted work improves or deteriorates.

Text and data mining exception

The copyright provision that would permit training on protected works. In the United Kingdom it remains narrow and unchanged.

Asset manifest

A seven-field record attached to an AI-touched asset at creation, described in Chapter 11.

Digital replica

A recognisable synthetic reproduction of a specific performer, governed in the United States by the 2026 collective agreements.

Assistive and generative

The distinction most sector surveys fail to draw, and the reason published adoption figures are not comparable.

Post-edited machine translation

Machine output corrected by a professional translator. The subtitling workflow with the strongest published evidence.

Process-based label

A disclosure describing how content was made, as distinct from a harm-based label describing whether it is misleading.

Errors and omissions insurance

Production insurance against claims arising from the content itself, and one of the gates at which AI provenance questions surface.

END MATTER

Method and evidence grading

How the sources in this paper were selected, checked and graded, and what the reader should distrust.

Sources were gathered from the peer-reviewed literature, primary legislation, court judgments, regulator and funder guidance, collective agreements and commissioned sector studies, with priority given to preregistered experiments, systematic reviews and primary official material over single studies and over vendor material, of which none is cited. Every reference in the list that follows was taken from the publishing journal, publisher, court or issuing body itself, and the position is the position as at 11 September 2026.

What could not be verified

Four things were searched for and not found, and the pattern of absence shapes this paper more than any single finding. No peer-reviewed study with a baseline and comparison for artificial intelligence in visual effects, rotoscoping, restoration, previsualisation or virtual production. No study measuring audience comprehension or reception of post-edited machine-translated subtitles against human-translated ones. No United Kingdom survey of artificial intelligence and screen employment with a published method and sample. And no concluded United Kingdom collective agreement on artificial intelligence between producers and performers, writers or crew, which is an absence confirmed as far as a reasonable search allows, and no more than that.

Figures corrected before publication

Two figures in common circulation did not survive checking and are not used here. The claim that image generation uses up to 1,000 times more energy than text tasks does not appear in the underlying research, which reports over 60 times more energy than text generation and reaches the larger multiple only by comparing a generative task with a non-generative one.¹⁴ And a single adoption rate for the United Kingdom screen sector does not exist, for the reasons Chapter 2 sets out.

What to distrust in this paper

Four things. The clearance cliff in Chapter 1, the consent scope in Chapter 7, the disclosure ladder in Chapter 9 and the asset manifest in Chapter 11 are the author's own frameworks, offered and not yet validated, and each is tagged as such wherever it appears. None has been tested outside the engagements it came from. The reading of Article 50 in Chapter 9 and of the copyright position in Chapter 6 is a summary of published sources and not a legal opinion. The transfer of the jagged frontier finding from consulting to screen craft is an argument, and it

carries whatever weight the reader thinks that argument deserves. And the workforce forecasts in Chapter 8 are scenario models commissioned by parties with an interest in the result, which the chapter states and the reader should weigh.

Conflict of interest

The author advises organisations on the deployment of artificial intelligence, including in regulated and creative settings, and would be a plausible supplier of several things this paper recommends. It also argues against several things that would be commercially straightforward to sell, including generative elements in delivered picture and any product positioned as a provenance guarantee. Readers should weigh the argument accordingly.

END MATTER

About the author

Paul Forrest is the founder of Ecaveo and works as a fractional head of AI with private equity and venture capital portfolio businesses on the deployment of artificial intelligence in regulated and rights-heavy settings. He began in business process reengineering in the 1990s and moved into natural language and machine learning work from 2014, building small, domain-specific models in preference to general-purpose ones.

He is the author of *Tales from AI Development Hell* and *AI Made the Decision: Who Answers When Nobody Owns the Outcome?*, and holds an adjunct professorship teaching digital acumen and digital entrepreneurship. A third book on where advantage accumulates in the use of artificial intelligence is in preparation.

This is the second paper in the Ecaveo AI adoption series. Comments and corrections are welcome.

END MATTER

References

43 sources, cited as at 11 September 2026.

Evidence on capability, quality and audience

1. Dell'Acqua, F., McFowland III, E., Mollick, E., Lifshitz-Assaf, H., Kellogg, K., Rajendran, S., Kraye, L., Candelon, F., & Lakhani, K. (2025). Navigating the jagged technological frontier. Field experimental evidence of the effects of artificial intelligence on knowledge worker productivity and quality. *Organization Science*. doi:10.1287/orsc.2025.21838. Preregistered field experiment, 758 consultants.
2. British Film Institute and CoSTAR Foresight Lab (2025). *AI in the Screen Sector. Perspectives and Paths Forward*. 9 June 2025. doi:10.5281/zenodo.15601301. Authors Angus Finney, Brian Tarran and Rishi Coupland.
3. Pact (2024). *Independent Production Sector Financial Census and Survey 2024*, conducted by Oliver and Ohlbaum Associates, with AI annex. Cited via Pact written evidence to Parliament, ACT0005, September 2024.
4. YouGov (2025). *AI in streaming entertainment. UK audiences want assistance, but not AI-generated content*. Fieldwork 23 to 25 October 2025, 2,035 adults aged 16 and over, Great Britain, nationally representative and weighted.
5. YouGov (2024). *Labelling AI-generated and digitally altered content*. Fieldwork 1 to 2 May 2024, 2,128 United Kingdom adults, nationally representative and weighted.
6. YouGov (2025). *Trust or trepidation. Attitudes to AI in media*. 27 February 2025. Fieldwork 16 December 2024 to 2 January 2025, 2,000 respondents, Great Britain.
7. Simon, F., Nielsen, R. K., & Fletcher, R. (2025). *Generative AI and news report 2025*. Reuters Institute for the Study of Journalism, 7 October 2025. Approximately 2,000 respondents in each of six countries.
8. Centre National du Cinéma et de l'Image Animée (2024). *Observatoire de l'IA. Perception par le public de l'utilisation de l'IA dans la création*. 11 September 2024. Fieldwork May 2024. Sample size not published.
9. Bywood, L., Georgakopoulou, P., & Etchegoyhen, T. (2017). Embracing the threat. Machine translation as a solution for subtitling. *Perspectives. Studies in Translation Theory and Practice*, 25(3), 492-508. doi:10.1080/0907676X.2017.1291695. SUMAT project; 37,104 subtitles post-edited by 19 professional subtitlers.
10. Sandoval-Castañeda, M., Copti, S., & Shasha, D. (2024). AutoTag. Automated metadata tagging for film post-production. *Multimedia Tools and Applications*, 83(3), 6731-6753. doi:10.1007/s11042-023-15565-w
11. Lavan, N., Irvine, M., Rosi, V., & McGettigan, C. (2025). Voice clones sound realistic but not (yet) hyperrealistic. *PLOS ONE*, 20(9), e0332692. doi:10.1371/journal.pone.0332692
12. Bakkouche, L., McGhee, C., Lau, E., Cooper, S., Luo, X., Rees, M., Alter, K., Post, B., & Schwarz, J. (2025). Finding the human voice in AI. Insights on the perception of AI-voice clones from naturalness and similarity ratings. *Interspeech 2025*, Rotterdam.
13. Graham, C., & Roll, N. (2024). Evaluating OpenAI's Whisper ASR. Performance analysis across diverse accents and speaker traits. *JASA Express Letters*, 4(2), 025206.
14. Luccioni, A. S., Jernite, Y., & Strubell, E. (2024). Power hungry processing. Watts driving the cost of AI deployment? *ACM Conference on Fairness, Accountability and Transparency (FAcT '24)*. arXiv:2311.16863. 88 models, 10 tasks, 1,000 inferences per model.
15. Altay, S., & Gilardi, F. (2024). People are skeptical of headlines labeled as AI-generated, even if true or human-made, because they assume full AI automation. *PNAS Nexus*, 3(10), pgae403. doi:10.1093/pnasnexus/pgae403. Two preregistered experiments, 4,976 participants.
16. Wittenberg, C., Epstein, Z., Péloquin-Skulski, G., Berinsky, A. J., & Rand, D. G. (2025). Labeling AI-generated media online. *PNAS Nexus*, 4(6). doi:10.1093/pnasnexus/pgaf170. Two preregistered experiments, 7,579 participants.

Law, policy and litigation

17. Department for Science, Innovation and Technology, Department for Culture, Media and Sport, and Intellectual Property Office (2026). *Report on Copyright and Artificial Intelligence*, CP2602959, and *Copyright and Artificial Intelligence. Impact Assessment*. Both published 18 March 2026.
18. Data (Use and Access) Act 2025 (c. 18), sections 135 and 136. Royal Assent 19 June 2025. Economic impact assessment and report on the use of copyright works in AI development.
19. Regulation (EU) 2024/1689 (the AI Act), Article 50. Transparency obligations for providers and deployers of certain AI systems, applying from 2 August 2026, including the artistic and creative works provision at Article 50(4).
20. *Bartz v Anthropic PBC*, Northern District of California, Judge William Alsup, 23 June 2025. Training held transformative; acquisition of pirated works held not fair use. Subsequently settled.

21. **Kadrey v Meta Platforms**, Northern District of California, Judge Vince Chhabria, 25 June 2025. Training on copyrighted books held fair use on the evidence presented. Proceedings continuing.
22. **Getty Images v Stability AI**, High Court of England and Wales, Mrs Justice Joanna Smith, 4 November 2025. Primary copyright claims withdrawn for want of evidence of copying in the United Kingdom; secondary infringement failed; trade mark claim succeeded.
23. **Thomson Reuters v ROSS Intelligence**, District of Delaware, Judge Stephanos Bibas, 11 February 2025. Fair use defence failed as a matter of law. Under appeal to the Third Circuit.
24. **Regulation (EU) 2026/1744** (the Digital Omnibus on AI). In force 27 July 2026. Postponed high-risk obligations; did not delay Article 50; transitional marking period to 2 December 2026 for systems placed on the market before 2 August 2026.
25. **European Commission** (2026). *Code of Practice on Transparency of AI-generated Content*, final version 10 June 2026, and *Guidelines on transparency obligations for providers and deployers of certain AI systems*, 20 July 2026.
26. **British Film Institute** (2026). *Guidance on the use of Artificial Intelligence (AI) in BFI funding applications and funded projects*. Last updated 28 January 2026.

Agreements, unions and sector bodies

27. **SAG-AFTRA** (2026). *2026 TV/Theatrical Contracts*. Ratified 4 June 2026 by 91.42 per cent of those voting. Term 1 July 2026 to 30 June 2030. Provisions on digital replicas, synthetic performers, scanning and biometric data.
28. **Writers Guild of America** (2023). *Minimum Basic Agreement*, artificial intelligence provisions, as summarised in the Guild's own Know Your Rights guidance.
29. **Equity** (2023, updated). *AI Toolkit* and *AI Vision Statement*, developed as part of the Stop AI Stealing the Show campaign, and the General Secretary's statement of June 2026 on the SAG-AFTRA standard as a United Kingdom benchmark.
30. **Equity** (2025). Indicative ballot on artificial intelligence protections for performers, results announced 18 December 2025. 99.6 per cent in favour on a turnout of 75.1 per cent of 7,732 eligible film and television performers.
31. **CVL Economics** (2024). *Future Unscripted. The Impact of Generative Artificial Intelligence on Entertainment Industry Jobs*. 12 February 2024. Survey of 300 industry leaders, fieldwork 17 November to 22 December 2023. Commissioned by the Animation Guild, the Concept Art Association, the Human Artistry Campaign and the National Cartoonist Society Foundation.
32. **CISAC and PMP Strategy** (2024). *Study on the economic impact of Generative AI in the Music and Audiovisual industries*. December 2024. Expert elicitation and forecast modelling to 2028. Commissioned by CISAC.
33. **Bectu** (2025). *Big Survey*. 5,597 respondents across the United Kingdom creative industries, more than 3,600 from television. The survey does not address artificial intelligence.
34. **BAFTA albert** (2025). *Carbon Calculator Methodology paper*. The methodology contains no treatment of artificial intelligence, cloud computing or data centre emissions.

Standards, provenance and security

35. **Coalition for Content Provenance and Authenticity** (2026). *Content Credentials. C2PA Technical Specification*, version 2.4, April 2026.
36. **Kaye, K., & Dixon, P.** (2025). *Privacy, Identity and Trust in C2PA. A Technical Review and Analysis of the C2PA Digital Media Provenance Framework*. World Privacy Forum, 3 September 2025.
37. **Ofcom** (2025). *Deepfake Defences 2. The Attribution Toolkit*. 11 July 2025. Includes an original technical evaluation of three watermarking tools across 26 edit types.
38. **Zhao, X., Zhang, K., Su, Z., Vasan, S., Grishchenko, I., Kruegel, C., Vigna, G., Wang, Y.-X., & Li, L.** (2024). Invisible image watermarks are provably removable using generative AI. *NeurIPS 2024*. arXiv:2306.01953
39. **Ofcom** (2026). *Ofcom's strategic approach to AI, 2026/27. Enabling safe and secure AI adoption*. 4 June 2026.
40. **Advertising Standards Authority and Committee of Advertising Practice** (2025). *Disclosure of AI in Advertising. Striking the Balance Between Creativity and Responsibility*. 29 May 2025.
41. **National Cyber Security Centre** (2023). *Guidelines for secure AI system development*. 27 November 2023.
42. **National Institute of Standards and Technology** (2024). *Artificial intelligence risk management framework. Generative artificial intelligence profile*, NIST AI 600-1. 26 July 2024. doi:10.6028/NIST.AI.600-1. The document is unrevised; the web record carries a later last-updated date.
43. **Information Commissioner's Office**. *Guidance on AI and data protection*, last updated 15 March 2023 and under review following the Data (Use and Access) Act, and *Biometric data guidance. Biometric recognition*, 13 November 2024.

A production intelligence design workshop

The argument in this paper is testable, and the cheapest way to test it is to take one delivered asset from a recent title and try to establish what it is made of. Most organisations cannot, and the exercise takes an afternoon. It changes the conversation more than any demonstration.

Ecaveo runs design workshops with studios, production companies, post facilities, funders and broadcasters. The output is a manifest the organisation owns, a consent scope drafted for the tools that exist, a clearance tier for each use in the estate and one evaluation protocol ready to run. It is not a product demonstration.

What a workshop involves

- Two sessions with the people who actually make the assets, not only with the people who buy the software.
- One delivered asset traced backwards, to establish what can and cannot currently be proved about it.
- A manifest specification that fits the asset management system the organisation already runs.
- A consent scope covering purpose, term, territory, medium and transfer, read against the current bargained floor.
- One pre-registered evaluation protocol for a named task class, with its stop conditions agreed in advance.

