

Findable

A Biometric Measurement of Intimate Imagery
Beyond the Reach of Ordinary Search

QayemDigital Research

Correspondence: research@qayemdigital.com

Version 1.4, 12 August 2026

Independent preprint; not peer reviewed

Competing interest. QayemDigital develops face-search technology. The study’s authors therefore have a commercial interest in the problem of biometric discovery described here. The study does not evaluate or endorse a QayemDigital product.

Cite as. QayemDigital Research, “Findable: A Biometric Measurement of Intimate Imagery Beyond the Reach of Ordinary Search,” version 1.4, 12 August 2026. Published by QayemDigital Research at <https://qayemdigital.com/blog-findable.html>, which is the version of record.

Abstract

Content moderation for non-consensual intimate imagery (NCII) is report-driven: material is removed when someone reports it. A report presupposes that the depicted person can *find* the material, and that prerequisite has, to our knowledge, never been measured. We measure it directly across four adult platforms. First, in 3,950 sampled clips (uniform random draws on three platforms, an uploader walk on the fourth), about half carry no name, alias, or production identifier in their metadata: no searchable person or production identifier is available. The locked classification pipeline finds 48.1% among the three uniformly sampled platforms and 56.1% in the observed four-platform sample after adding erome’s non-uniform uploader walk; a model-blind reviewer-led audit ($n = 400$) and a recorded exact-model cross-provider pass bound the same quantities at 40.5–48.1% and 49.2–56.1%, so we report “about half” rather than a settled majority. Second, biometric face clustering of 2,497 encountered uploader catalogues, at a threshold calibrated on performer-labelled adult frames, shows that **90% of the 1,943 accounts crawled to assessable depth (≥ 10 posts) depict ten or more biometrically distinct people**; an inverse-multiplicity reweighting leaves the figure essentially unchanged (89.5%). Together the measurements identify a *discovery gap*: metadata search and account-by-account monitoring are unreliable routes to the measured content, so faster takedown cannot reach items that are never found. The claims concern discoverability, not consent, authorship, or total prevalence. The metadata measure does not establish that the depicted person knows the term, that the term is sufficiently specific, or that a search engine retrieves the clip. This is not an estimate of NCII prevalence; it measures discoverability-related properties of intimate imagery in the sampled platform content.

1 Introduction

Adult platforms remove non-consensual intimate imagery (NCII) and leaked material principally through notice-and-takedown: a report is filed, and the platform acts on it. Statutory removal

duties enacted across a growing number of jurisdictions, including the TAKE IT DOWN Act in the United States, the Online Safety Act in the United Kingdom, and the Digital Services Act in the European Union [1, 2, 3], differ in who may notify a service and what duties apply. They nevertheless operate alongside reactive mechanisms that require particular content to be located before it can be notified. Reporting is therefore an important, but not exclusive, route into the moderation regime.

That reactive route rests on an under-measured premise. Before anything can be reported it must first be *discovered*: the depicted person, or someone acting for them, must locate the material. To our knowledge, prior work has not quantified whether discovery is possible. The literature on image-based sexual abuse quantifies harm that was reported or disclosed; it does not quantify the share of platform content carrying a searchable person or production identifier. When discovery fails, faster takedown cannot help that particular item because no report is filed.

This paper measures discoverability directly. We ask two empirical questions:

1. **Searchability.** Can the depicted person find a given clip? Does its metadata contain anything they could search?
2. **Concentration.** How is imagery organised across upload accounts? Are accounts centred on one person, or are they collections of many?

We answer the first with a sample of clips classified for the presence of an searchable person or production identifier, and the second with biometric face clustering of uploader catalogues. We use *discovery gap* for the measured weakness of two ordinary routes: metadata search and monitoring a single uploader page. We do not treat those measurements as a census of every possible discovery route.

We frame the study as the measurement of a privacy failure, namely how little of this imagery ordinary search can surface for the person depicted, and not as a study of copyright or piracy. Throughout, “intimate imagery” denotes a single category. We do not sub-classify by origin or consent status, neither of which is observable at scale, and our claims do not depend on either.

2 Background and Related Work

The report-driven model. Removal of NCII on these platforms often depends in practice on a reporter who identifies specific URLs. Platform report forms and DMCA notices take a located instance of the material as input, and the empirical notice-and-takedown literature likewise treats a located item as the unit of enforcement [13]. Statutory regimes differ, and the DSA’s notice mechanism, for example, is open to individuals and entities rather than victims alone [1, 2, 3]. The reactive part of the model assumes that some notifier first finds the item.

Harm measurement. The research literature on image-based sexual abuse has established the scale and severity of the harm through survey work, victim reports, and legal analysis [4, 5, 6, 7, 16, 17]. These sources quantify harm that was reported or disclosed; they do not estimate the platform-side quantity measured here, the share of sampled clips that carry a searchable person or production identifier.

Moderation infrastructure and proactive routes. Platform moderation is organised around queues of located items [14, 15]. Proactive mechanisms exist for *known* images: perceptual hash matching such as PhotoDNA [18], and victim-initiated hash enrolment through StopNCII [19] and,

for minors, NCMEC’s Take It Down service [20], can block re-uploads of an image the person already possesses. These routes presuppose a report or a source image in hand; none helps a person find material they have never seen.

Face search. Consumer face-search engines index faces on public pages [11, 21]. Their coverage of the platforms measured here is unknown, and biometric identification engages special-category data rules [12]; §5 returns to both points. Synthetic imagery further widens the pool of material a depicted person has never seen [22]; we do not sub-classify origin, and our claims do not depend on it.

Measurement of adult platforms. Large-scale measurement of the adult web has characterised its tracking and privacy properties [23]. We are not aware of prior work quantifying this metadata searchability construct or per-account depicted-person counts on these platforms.

Why these platforms resist measurement. The properties that make these platforms hard to study are themselves part of the finding. Catalogue enumeration differs across sites, re-encoding and re-uploading can strip provenance, and, as we show, many sampled clips carry no name at all. The material is difficult to enumerate by design and by practice, which is why the discoverability question has remained open.

3 Data and Methods

Table 1 summarises the design. Each finding pairs a primary instrument with independent checks, and every correction made during the study is documented with receipts (§7, Appendix C).

Table 1: Design summary.

Component	Sample	Instrument and checks
Finding 1: searchability	3,950 clips: uniform random draws on epornr, tnaflx, spankbang ($n=1,000$ each); uploader walk on erome ($n=950$)	Locked codebook; Haiku + Sonnet with Opus adjudication; model-blind reviewer-led audit ($n=400$); recorded exact-model GPT-5.6 Sol cross-provider pass
Finding 2: concentration	2,497 encountered ordinary-user accounts, of which 1,943 assessable (≥ 10 crawled posts)	SCRFD + ArcFace clustering at cosine ≥ 0.42 , calibrated on 85 performers; 50-account re-fetch; 469-pair blinded candidate review; inverse-multiplicity reweighting

3.1 Platform selection

We measured four platforms with distinct architectures: **epornr** and **tnaflx** (large video tubes), **spankbang** (video tube, Cloudflare-gated), and **erome** (image and video galleries).

3.2 Sampling

Where a platform publishes its full catalogue we drew a uniform random sample: epornr via its sitemap index (approximately 2.25M videos), tnaflx and spankbang via random draws over their

numeric identifier space. Where no full catalogue is exposed (erome), we used an uploader walk, which yields a sample stratified by discovered uploader rather than a uniform one; this limitation is stated and carried through the analysis (§4.1, §6). We sampled approximately 1,000 clips per site (3,950 total; erome 950) in July 2026, logging the draw method and date for reproducibility.

The Finding 2 account frame was derived from those clips. Deduplicating their uploaders produced 2,804 accounts: 2,497 labelled by the platforms as ordinary users, 192 as studios, and 115 as verified creators. The 2,497 user accounts covered 3,557 of the sampled clips and were the catalogues sent to the crawl and clustering stage. This is an *account set encountered through a clip sample*, not an account-uniform probability sample. Accounts that publish more clips have more opportunities to enter the frame in the clip-uniform components, while erome inherits the uploader-walk design; after deduplication, the results weight each encountered account equally. Finding 2 therefore describes this encountered uploader set and must not be generalized as an unbiased prevalence estimate for all platform accounts.

3.3 Searchability classification

We locked the operational definition before analysis. A clip is **searchable** if its metadata contains either (a) a personal name, alias, or handle of the depicted person, or (b) a searchable production identifier, meaning a studio, scene, or film code (for example, catalogued adult-film codes). It is **unsearchable** otherwise. Generic descriptors (body type, act, category), dates, and quality tags do not count as identifiers. A full legal name or a globally unique alias is not required: a one-word alias counts when the metadata presents it as the depicted person’s name. Conversely, an exact but generic natural-language title is not a production identifier merely because the phrase could be searched; it must expose a recognisable studio, series, scene, film, or catalogue key. The measure concerns whether the metadata provides a searchable person or production identifier as a potential search seed, not guaranteed retrieval. It does not establish that the depicted person knows the term, that the term is sufficiently specific, or that a search engine retrieves the clip. Synthetic codebook examples appear in Appendix B; source identifiers are withheld.

Classification used three separate model runs. Claude Haiku and Claude Sonnet classified every clip’s metadata and agreed on 91.9% of cases; Claude Opus adjudicated the 318 disagreements. These are separate model families, not three independent observations in a statistical sense. The scripts preserve the exact Haiku identifier (`claude-haiku-4-5-20251001`); the exact Sonnet and Opus identifiers were not written into the result records, a reproducibility omission listed in §6. Classifiers were constrained to cite only names appearing verbatim in the metadata, which suppresses inferred or invented identifiers. Candidate production identifiers flagged by the first model were confirmed by the second before counting towards searchability; 410 of 612 candidates survived confirmation, and the 82 unconfirmed candidates on otherwise nameless clips are counted as unsearchable.

An initial name-detection audit and an initial full-searchability audit contained 100 items each, with fixed seeds 7 and 11 and deliberate over-sampling of hard cases. The name audit produced 74.0% raw and 83.5% stratum-weighted agreement. The initial searchability audit produced 73.5% raw and 77.6% stratum-weighted agreement, but was too small for a precise prevalence estimate; it is retained as provenance rather than used for the final reviewer-led estimate.

We therefore locked an expanded 400-item audit before collecting its decisions. The sample used fixed seed 20260724 and four historical-pipeline strata: 150 named controls, 50 confirmed-code records, 50 contested-code records, and 150 records with no detected identifier. Allocation within each stratum was proportional across platforms, with exact inclusion weights retained for every platform-by-stratum cell. The interface displayed only title, tags, uploader/channel, performer,

and description fields; it hid the historical verdict, stratum, platform, source link, and source key. It displayed no imagery. The reviewer made 391 binary decisions and nine abstentions. The interface also offered “name”, “production identifier”, and “both” as findable reasons, but those reasons were not applied consistently in the first 100 items. They are therefore collapsed to a single *findable* category and are not used for reason-specific prevalence claims or public release. Appendix D reports the expanded audit.

After the expanded audit, a provenance check found that the tnaflx uploader parser had captured a related-video uploader rather than the main video’s uploader on many pages. We preserved the July input, repaired the parser, re-fetched all 1,000 sampled tnaflx pages, and re-adjudicated the 100 affected audit records under new opaque IDs with prior decisions, platform, status, strata, and source keys hidden; 11 collapsed findable/unfindable/abstain decisions changed, and the merged corrected audit is the final audit reported below. The production name classifier used only title and tags, and a deterministic check confirmed that no positive name decision depended on the repaired uploader field. Appendix C details the repair, its receipts, and the matching re-run of the recorded exact-model cross-provider pass.

The expanded audit was reviewer-led, not human-only. For some linguistically opaque or codebook-edge records, the reviewer consulted an AI assistant or dictionary for translation and application of the written rules. The assistant was not shown the historical model verdicts, strata, platform labels, or linkage keys. During the original and correction reviews, source pages were consulted in a small number of opaque abbreviation or account-status cases; four such records were ultimately labelled findable. We report the conservative sensitivity to treating all four as unsearchable in §4.3. These qualifications prevent the audit from being described as an independent human-only benchmark, while preserving its independence from the historical pipeline’s displayed decisions.

Finally, we ran a cross-provider sensitivity pass over all 3,950 records with the exact requested and returned model identifier `gpt-5.6-sol`, medium reasoning, the Responses API, and `store:false`. The locked prompt, input, and restricted decision output have SHA-256 receipts. This pass is a new robustness analysis, not a reconstruction of the unrecorded historical Sonnet or Opus versions and not a substitute for reviewer adjudication. After the uploader repair, all 1,000 tnaflx records were re-run with the same recorded model ID, prompt, and settings; the other 2,950 outputs were retained. The merged corrected pass and its separate tnaflx provenance receipt are the cross-provider results reported below.

Some platform-labelled creator and studio uploads may be self-posted. We therefore report a deliberately broad sensitivity analysis that counts all such uploads as searchable, while recognising that account status does not prove the depicted person is the uploader (§4.1). Status was detected per platform: epornr and spankbang mark studio and creator uploads with a distinct account namespace, tnaflx carries a verified badge, and erome marks verified creators with a profile check.

3.4 Biometric instrument

Faces were detected with SCRFD [8], aligned to a 5-point template at 112×112 , and embedded with an ArcFace ResNet-100 model trained on Glint360K [9, 10] to a 512-dimensional unit vector; identity similarity is cosine similarity.

We calibrated the operating threshold in two stages. An initial benchmark on labelled public-figure images gave substantial separation (precision 1.000, recall 0.955 at cosine 0.38). Because clean images do not reflect the operating domain, we recalibrated on real, performer-labelled adult frames: 85 performers drawn automatically from the platform’s performer directory, with identity established independently of the matcher. On these frames, same-performer and different-performer

similarities show substantial separation (median cosine 0.44 versus 0.03). We adopted **cosine** ≥ 0.42 , at which 31 of 237,904 different-performer pairs match (a false-match rate of 1.3×10^{-4} , roughly half the rate at 0.38), with the residual false matches confined to similar-looking performers (Appendix A). Because comparisons are clustered within 85 performers, this is an empirical pair-level false-match rate, not an independent-observation confidence estimate. An earlier calibration on 25 performers (41,734 pairs) gave a materially identical figure (1.2×10^{-4}), so the threshold is not tuned to the size of the calibration set. As an independent check on the labelling, a reviewer who could personally identify two of the performers verified all 108 automatically labelled faces for them; 105 (97%) were correct.

For each account we clustered its faces by transitive graph expansion, linking any two faces at cosine ≥ 0.42 ; each connected component is one identity. The cluster count is the number of biometrically distinct people the account’s content depicts. We emphasise that this **measures depiction, not authorship**: the count reflects who appears in the content, not who operates the account, and a person who posts their own content may appear with partners.

Three properties of the clustering deserve explicit treatment. First, transitive expansion can in principle chain two distinct identities through a borderline intermediate face. The false-match rate at the operating threshold bounds each such link at roughly one in 10^4 pairs, and, more importantly, chaining *merges* clusters: it lowers the distinct-person count and therefore cannot inflate the ≥ 10 headline. The calibration’s residual failure mode, merging similar-looking performers, runs in the same conservative direction. Second, the error that could inflate counts is over-splitting (one person’s degraded frames landing in separate clusters). We designed a candidate-pair medoid review for that failure mode, but withdrew the initial qualitative result after provenance checks exposed contaminated pre-clean spankbang review material. In the completed repaired review, 11 confirmed same-person pairings reduced the repeat counts of nine candidate-bearing accounts by 11 identities in total (at most three in one account). No reviewed account changed, or could change if every uncertain pair were treated as the same person, its ≥ 10 classification (§3.5). This checks threshold robustness within the bounded candidate screen; it does not prove that no unselected over-split pair exists. Third, the calibration covers 85 performers, and false-match estimation operates on pairs, of which it yields 237,904 different-performer comparisons; its identity coverage is a stated limitation rather than a hidden one.

3.5 Repeatability, validity, and error analysis

We re-fetched a stratified 50-account sample (spanning low, mid, and high distinct-person counts) and re-clustered. In the repaired record, per-account recounts have Pearson $r = 0.997$, a median relative difference of 0%, and 48 of 50 accounts within 25% of the original count. Ten low-count accounts produced zero or one repeat medoid and therefore cannot test whether one person was split across clusters; the human comparison denominator is 40 accounts.

Preparing machine-readable human verdicts exposed a contamination defect in the validation set: exact crop hashes showed that its 15 spankbang rows predated the strip-excluded re-crawl. The affected rows and their qualitative review were withdrawn, archived, and deterministically replaced from the clean 742-account crawl (Appendix C). The correction concerns the validation sample only; the headline account table already used the separate clean re-crawl, and the production face index was not affected.

The re-fetch establishes computational repeatability, not identity validity. Pairwise calibration tests separation at the chosen threshold; a repaired candidate-pair review targets the principal count-inflating failure mode, over-splitting. We re-embedded all 1,654 stored face-only medoids with the same antelopev2/glintr100 recognizer. Within each account, a pair enters review if its cosine is at

least 0.30, its decoded RGB pixels are identical, or its 64-bit difference hash has Hamming distance at most 5. Detector-recovery passes achieved complete medoid coverage; their thresholds and transforms are recorded in the generator and private provenance file. The rule selected 469 pairs across 25 accounts. For scale, one 93-medoid account would require 4,278 exhaustive comparisons but yields nine candidates.

The restricted interface shows one pair at a time in deterministic blinded order, hides selection scores, and records *same*, *different*, or *uncertain* with pair-level confidence and private notes. Confirmed same-person edges are transitively grouped; the corrected count and ≥ 10 classification effect are derived automatically. Full-account medoids are optional context rather than the unit of review. The one reviewer adjudicated all 469 prespecified candidate pairs: 445 were *different*, 11 *same*, and 13 *uncertain*. Confirmed edges reduce counts by 11 identities across nine of the 25 candidate-bearing accounts, with a maximum correction of three in one account. Fifteen candidate-bearing accounts had only different decisions and six contained at least one uncertain decision. None of the 25 changed its ≥ 10 classification, and none could change if all uncertain pairs were treated as same-person edges.

An interface ambiguity was also corrected during collection: an early *Next* control had been used to mean *different* without persisting the choice. The affected decisions were reconstructed from deterministic request logs, a 34-position range was reset and re-adjudicated, and before/after archives with machine-readable receipts preserve the transition (Appendix C).

This is a reproducible, bounded screen, not exhaustive identity adjudication: over-split pairs below cosine 0.30 that also evade the two image-hash rules can be missed, and an account with no selected candidate is not thereby human-validated across every possible pair. Because selection deliberately enriches plausible duplicates, 11/469 is neither a population error rate nor a recall estimate. The restricted account-level record, repaired 50-account validation record, archived predecessor, candidate-pair provenance, de-identified verdict output, and correction receipts preserve the audit trail without publicly exposing source identifiers or face crops.

4 Results

4.1 Finding 1: about half lack a searchable person or production identifier

The per-site rates are the primary result. Across the three uniformly sampled platforms, **48.1%** of 3,000 clips have no searchable person or production identifier in their metadata. Adding erome’s non-uniform uploader-walk sample gives **56.1%** in the observed four-platform sample across 3,950 clips. Under the deliberately broad sensitivity assumption that all verified-creator uploads are searchable, that four-platform share falls to 54.2%; extending the assumption to all studio uploads gives 53.0%. Account status does not prove who is depicted, so these are sensitivity bounds rather than observed self-post rates. The per-site pattern (Table 2) is a gradient: epornr is lowest and erome is highest. This is consistent with professionally catalogued content carrying names and production codes more often, but the study does not test that explanation causally.

Neither pooled figure is an account- or catalogue-weighted platform prevalence estimate. The 48.1% figure holds the sampling design to uniform clip draws but excludes erome; the 56.1% figure includes all four measured platforms but is raised by erome’s uploader-walk sample. We therefore use “about half” as the cross-site summary and do not treat either pooled value as establishing a general majority.

The two added checks place the automated result in a narrower measurement range (Table 3). The expanded audit estimates **52.6%** unsearchable across all four samples (95% stratified normal CI 49.7–55.4) and **43.9%** across the three uniform samples (40.7–47.2). Survey-weighted reviewer–

Table 2: Finding 1: automated share of clips with no searchable person or production identifier in metadata. The pooled parenthetical is the sensitivity range for the treatment of likely self-posted content (§4.3); the descriptive 95% Wilson interval on the unadjusted pooled rate is 54.5% to 57.6% and is not a platform-population sampling interval.

Platform	n	Unsearchable
eporner	1000	37.4%
spankbang	1000	52.9%
tnaflix	1000	54.1%
erome	950	81.2%
Uniform draws (3 sites)	3000	48.1%
Pooled (sampled)	3950	56.1% (53.0–56.1)

pipeline agreement is 91.3% with Cohen’s $\kappa = 0.825$. The recorded exact-model GPT-5.6 Sol pass finds 49.2% and 40.5%, with 92.4% agreement and $\kappa = 0.847$ against the historical pipeline across all 3,950 records. These comparisons preserve the platform gradient and the “about half” interpretation, but show that the exact rate is instrument-sensitive.

Table 3: Finding 1 triangulation. Audit intervals are stratified normal 95% confidence intervals over decided items; GPT-5.6 Sol is a recorded exact-model cross-provider sensitivity pass on the same records, not a human-validity estimate.

Instrument	All four samples	Uniform three platforms
Historical Claude pipeline	56.1%	48.1%
Reviewer-led audit ($n = 400$)	52.6% (49.7–55.4)	43.9% (40.7–47.2)
GPT-5.6 Sol robustness pass	49.2%	40.5%

Because searchable production codes were counted on the *searchable* side, this measures the residual layer of content with neither a name nor a confirmed code: the layer for which no searchable person or production identifier is available in the metadata. Within that layer, 82 clips (2.1% of the sample) carried a candidate code that failed confirmation. In the expanded audit, 39.6% of the contested-code stratum remained unsearchable and raw reviewer–pipeline agreement there was 39.6%, making it the least stable measurement boundary. The cross-provider comparison points to the same issue: name detection agreed on 97.6% of records, while GPT-5.6 Sol accepted 879 production identifiers compared with 410 in the historical pipeline. This is a codebook boundary, not evidence that either model is correct (Appendix D). Appendix B records worked boundary examples. We decline to widen exposure by adding raters, even for metadata-only review (§7); the identified next steps for this boundary are further codebook hardening and a blind intra-rater re-test by the same reviewer under fresh opaque identifiers.

4.2 Finding 2: encountered uploader accounts depict many people

We clustered the catalogues of the 2,497 ordinary-user upload accounts encountered through the Finding 1 clip sample. Among those crawled to an assessable depth (≥ 10 posts), the study reports that **94% on the three deep-catalogue platforms, and 90% across all four**, depict ten or more biometrically distinct people (Table 4). Clustering is within-account: these counts say how many people appear in one account’s catalogue, not how many accounts any one person appears in

(§6). The reported median assessable account’s content shows a different person in roughly half of its posts (median diversity ratio, distinct people divided by posts, of 0.49). Approximately 9% of accounts reportedly depict a single distinct person, which is consistent with self-posting but does not establish it.

Table 4: Finding 2: reported distinct people depicted per encountered ordinary-user upload account. “ $\% \geq 10$ ” is conditioned on accounts with ≥ 10 crawled posts; assessable accounts number 1,943 of 2,497 in total, of which spankbang contributes 383 and the three deep-catalogue platforms 1,560. This is not an account-uniform sample.

Platform	n	Med. posts	Med. distinct	Diversity	$\% \geq 10$
eporner	471	185	81	0.53	95%
tnaffix	812	51	22	0.45	96%
erome	472	154	80	0.53	91%
spankbang	742	10	5	0.50	74%
Pooled (4 sites, 1,943 assessable)				0.49	90%
Three deep-catalogue sites (1,560 assessable)					94%

In the repaired 50-account validation sample, 25 accounts produced candidate pairs. All 469 prespecified candidates were adjudicated. Confirmed same-person decisions reduced counts by 11 identities across nine accounts, with a maximum reduction of three; 13 pair decisions were uncertain. No candidate-bearing account crossed the ≥ 10 threshold, and none could cross if every uncertain pair were treated as the same person. This supports the stability of the headline classification in this validation sample under the bounded screen, not the absence of clustering error.

We report both the four-site (90%) and three-site (94%) figures. Platforms were crawled to very different depths; median posts crawled ranged from 10 to 185. An account cannot depict ten distinct people if fewer than ten of its posts were sampled, so we condition on an assessable catalogue. Of the four platforms, spankbang was reachable to the shallowest depth because it rate-limits most aggressively, and its lower figure is partly, though not wholly, a crawl-depth artefact: normalised by depth, its diversity ratio (0.50) is close to the other platforms, though its accounts do skew smaller and more often single-person (30%).

Because the account frame is encountered rather than uniform, accounts that contributed more sampled clips have higher inclusion multiplicity. As a sensitivity, we reweighted each assessable account by the inverse of its sampled-clip multiplicity (67.5% of assessable accounts entered through a single sampled clip; the maximum multiplicity is 40). The reweighted ≥ 10 share is 89.5% pooled, versus 90.3% unweighted, and 94.3% on the three deep-catalogue platforms, versus 94.2%; no per-site figure moves by more than 1.5 points (eporner, 94.8% to 93.3%). This adjusts only unequal encounter multiplicity within the sample; it cannot restore the unknown inclusion probabilities of never-encountered accounts, so it is a partial check rather than an account-uniform estimate. Its stability indicates that the headline is not driven by a few repeatedly encountered prolific accounts.

The interpretation is bounded by the instrument. The counts concern depiction and concentration only: they say how many people appear in an account’s content, not who uploaded it, and, because clustering was performed within accounts, they do not say how many accounts any one person appears in. Within the encountered uploader set, the reported counts indicate that many accounts depict many different people rather than being organised around one person. The observed concentration is inconsistent with a catalogue consisting predominantly of material featuring the account operator alone and is compatible with an account-level collection depicting many different

people. The study does not establish how that material was obtained, uploaded, or distributed.

4.3 Uncertainty and sensitivity

Descriptive binomial uncertainty within the observed clip samples is narrow. For the non-uniform four-platform observed sample, the 56.1% share has a descriptive Wilson interval of 54.5% to 57.6%; this is not a platform-population sampling interval. Per-site descriptive intervals in Table 2 are within ± 3.1 points at n of 950 to 1,000. The remaining uncertainty is dominated by measurement and sampling-frame limitations, including definitional sensitivity and classification error.

The definitional part is the treatment of self-posted content, which moves the historical pipeline rate between 53.0% and 56.1% (§4.1). For the classification part, the expanded audit estimates 52.6% unsearchable (95% CI 49.7–55.4) across all four samples and 43.9% (40.7–47.2) across the three uniform samples. These intervals use the locked platform-by-stratum weights and treat the nine abstentions as missing at random within their sampling cells. Assigning every abstention to the findable or unfindable side gives an all-four range of 51.3–53.6% and a uniform-three range of 42.2–45.3%. Four source-assisted records were labelled findable; a conservative metadata-only sensitivity that relabels all four unsearchable moves the estimates from 52.6% to 53.1% overall and from 43.9% to 44.6% on the uniform three.

The cross-provider pass estimates 49.2% overall and 40.5% on the uniform three. Its overall agreement with the historical pipeline is 92.4% ($\kappa = 0.847$), while the reviewer-led audit has 91.3% survey-weighted agreement ($\kappa = 0.825$). The historical, reviewer-led, and cross-provider all-four estimates therefore span 49.2–56.1%; on the uniform three they span 40.5–48.1%. The expanded audit’s lower confidence limit remains just below 50%, so it does not independently establish a majority. The convergence supports a substantial discovery gap reasonably summarised as “about half”, while the production-identifier disagreement prevents the exact percentage from being treated as instrument-free.

For Finding 2, descriptive binomial intervals conditional on the encountered assessable accounts are similarly narrow: 90% of 1,943 assessable accounts gives a 95% interval of 88.6% to 91.3%, and 94% of 1,560 gives 92.7% to 95.1%. There too the dominant uncertainty is the instrument rather than the sample, and its two error modes push in opposite directions: false merges, including transitive chains, lower distinct-person counts and cannot inflate the ≥ 10 share (§3.4), while over-splitting could inflate it. The completed repaired review found limited count corrections (11 identities across nine candidate-bearing accounts) but no actual or possible ≥ 10 classification change under its candidate rules (§3.5). These intervals still do not account for the non-uniform account frame beyond the multiplicity reweighting sensitivity (§4.2), unselected low-similarity over-splits, or single-reviewer uncertainty.

5 Interpreting the Discovery Gap

The two findings identify weaknesses in two ordinary discovery routes.

1. Reactive reports presuppose discovery: one cannot report a specific item without first locating it.
2. **Finding 1** weakens metadata search. Across the historical pipeline, expanded reviewer-led audit, and recorded exact-model cross-provider pass, the all-four estimate ranges from 49.2% to 56.1% and the uniform-three estimate from 40.5% to 48.1%, with the sampling and measurement qualifications above.

3. **Finding 2** weakens the idea that a person can monitor one natural uploader page. In the encountered account set, the reported biometric counts show that many uploader catalogues depict many people.
4. The study does not measure incidental discovery, reverse-image search, consumer face-search coverage, hash matching [18, 19], or cross-account recurrence.

The narrow conclusion is therefore about the two routes measured, not about every conceivable path to an item. Material can surface when an acquaintance recognises a person, when the depicted person already has a source image for reverse-image search, through a consumer face-search engine, or through a platform’s proactive detection. Those routes are real and were not evaluated here. The measurement shows why a notice-only workflow can miss material upstream of takedown; it does not establish what share is never discovered by any means.

Nor does the study establish that existing face-search products fail to find this content. PimEyes’ own documentation says it checks publicly accessible pages, including pages with explicit content, and that explicit results are hidden when a user chooses its Safe Search option [11]; we did not measure its coverage of the four platforms. Any future biometric discovery tool would require a separate legal and governance analysis. Processing biometric data for recognition can engage special-category data rules and requires an Article 9 condition in addition to an Article 6 lawful basis under UK GDPR guidance [12]. Whether self-search, authorised-agent search, consent, or another design satisfies those requirements depends on jurisdiction and implementation; this measurement does not declare a lawful basis.

6 Limitations

- **Sampling.** The erome sample is stratified by uploader, not uniform; no full erome catalogue is exposed. The pooled Finding 1 rate also weights platforms by sample rather than catalogue size; per-site rates are primary, and the three uniformly sampled platforms give 48.1% historically, 43.9% in the reviewer-led audit, and 40.5% in the cross-provider pass.
- **Finding 2 account frame.** The 2,497 ordinary-user accounts are uploaders encountered through the clip sample and then deduplicated. This is not a uniform sample of platform accounts; account-level percentages apply to that encountered set. A genuinely account-uniform estimator cannot be reconstructed from these records because none of the four platforms exposes a complete, enumerable ordinary-user account frame for the study window; resampling the encountered accounts would not repair the original inclusion probabilities. An inverse-multiplicity reweighting leaves the estimates essentially unchanged (§4.2), but it corrects only encounter multiplicity, not the unknown inclusion probabilities of never-encountered accounts.
- **Searchability is a proxy.** We measure whether the metadata provides a searchable person or production identifier. The measure does not establish that the depicted person knows the term, that the term is sufficiently specific, or that a search engine retrieves the clip. The production-identifier boundary is materially less stable than the name boundary across instruments.
- **Biometrics measure depiction, not authorship.** Distinct-person counts say who appears, not who uploaded; verified and self-post labels are the platform’s own. Clustering was within-account, so no per-person count of accounts exists.

- **Clustering error.** Counts are bounded by the clustering. Over-splitting on degraded frames can inflate counts. The initial human conclusion was withdrawn after contaminated validation material was found; the repaired candidate-pair review found 11 confirmed count reductions across nine accounts but no actual or possible ≥ 10 threshold change. Its bounded candidate rule can still miss low-similarity over-splits, and its enriched 11/469 pair result is not a population error rate.
- **Single reviewer and assisted edge cases.** Exposure minimisation shaped the validation design. One trained reviewer saw only the metadata and face-only material needed for the audits; adding reviewers would have widened access to explicit identifiers and biometric crops derived from sensitive content. This safeguard precludes inter-rater reliability. For some ambiguous metadata the reviewer used an AI assistant or dictionary for translation and code-book application, and source pages were checked in a small number of opaque abbreviation or account-status cases. A conservative sensitivity treats all four source-assisted findable records as unsearchable. The expanded audit is therefore reviewer-led, not a human-only benchmark. Its reason subcategories were also applied inconsistently in the first 100 items, so only the collapsed findable/unfindable decision is analysed and released (§4.3).
- **Restricted and distributed replication materials.** The reviewer-audit decisions, account-level Finding 2 rows, and 50-account validation record exist, but are stored across restricted study systems and are not suitable for open release because they contain or can be linked to explicit names, uploader handles, URLs, and biometric review material. The repaired pair annotation is complete, and its de-identified verdict output and aggregate completion receipt can be released without those linkage fields.
- **Model version provenance.** The exact Haiku identifier is recorded, but the Sonnet and Opus runs were supplied through runtime flags that were not preserved in outputs, logs, or shell history. Their model families and July 2026 run window are known; their exact minor versions are not. The new robustness pass closes this gap only for itself: `gpt-5.6-sol` was both requested and returned on every record, with prompt, input, and output hashes retained.
- **Unmeasured discovery routes.** We measure metadata search and account concentration; the coverage of consumer face search is not measured. Incidental discovery, notification by acquaintances, and reverse-image search from a known source image were not measured; our claims concern dependable mechanisms, not every conceivable path.
- **Calibration breadth.** The threshold calibration covers 85 performers (237,904 different-performer pairs). Its known residual error, look-alike merging, lowers distinct-person counts and is therefore conservative for Finding 2, but broader identity coverage would strengthen it further.
- **Crawl depth** varies across platforms and is controlled by conditioning on assessable catalogues; spankbang remains the shallowest.
- **Snapshot.** Measurements reflect the platforms in July 2026; catalogues change.

7 Ethics, Legal Basis, and Governance

7.1 Framework and oversight

QayemDigital Research is an independent organisation without access to an institutional review board. In place of institutional review, the study was governed by a written protocol with definitions locked before analysis, and this section is structured around the Menlo Report principles for information and communication technology research [24] (respect for persons, beneficence, justice, and respect for law and public interest), so the reasoning itself can be examined. We treat the absence of independent ethical oversight as a limitation of the work, and we invite correspondence from ethics reviewers (research@qayemdigital.com).

7.2 Respect for persons and minimisation

Individual consent is impossible at this scale, and seeking it would itself require identifying depicted people, which the design forbids. The study therefore measures a privacy harm under strict minimisation. All reported figures are aggregate; **the paper identifies, ranks, or links no sampled individual to an uploader or real-world identity**. No intimate scenes or full frames are published. Calibration and quality control used restricted face-only crops and biometric templates; performer labels were used only to evaluate the threshold, and study accounts were not cross-searched against named identities. Human review used face-only candidate pairs and optional medoid context rather than intimate scenes.

7.3 Beneficence and justice

The processing burden falls on people depicted in the sampled content, including, unavoidably, the population the study aims to serve. The benefit is a quantity that cannot otherwise be known: whether the reporting regime’s implicit premise, that depicted people can find their material, actually holds. The risk concentrates in re-identification, which the design controls by publishing no identifiers, no imagery, and no linkage material, and by restricting the underlying records (see Data and Code Availability). In our judgement, aggregate measurement under these controls adds no exposure for any depicted person beyond the platforms’ existing publication of the content.

7.4 Crawling posture

Measurements used only publicly served pages, fetched with browser-equivalent HTTP clients. No accounts were created, no logins were performed, no paywall or age-verification mechanism was bypassed, and no private APIs were accessed. Requests were paced, with backoff on rate limiting; where a platform rate-limited aggressively (spankbang), we accepted the resulting shallower crawl rather than evading it (§4.2).

7.5 Special-category data

Face embeddings processed for identification are special-category biometric data under UK and EU data protection law [12]. To the extent that law applies to this processing, the research-purposes condition with appropriate safeguards is a candidate Article 9 basis: the purpose is scientific measurement in the public interest, and the controls include minimisation, no publication of biometric material, no decision or effect concerning any individual, restricted storage, and retention limits. We do not rely on any claim that depicted persons made their own data public, which would assume the self-posting this study does not assume. This is not a complete legal determination. UK GDPR

guidance requires a separate Article 6 lawful basis as well as an Article 9 condition [12]; territorial scope and jurisdictional variation also require analysis. This paper offers no external legal opinion and does not claim that its research-purpose discussion, by itself, supplies the necessary Article 6 analysis or a formal data-protection impact assessment. A documented impact assessment and jurisdiction-specific legal review are required before any future continuation of the biometric processing or any product deployment. The measurement does not declare a lawful basis for any other actor’s processing (§5).

7.6 Retention and deletion

Sensitive artifacts (face crops, embeddings, medoid indexes, montages, raw metadata, and linkage keys) remain on restricted systems. They will be deleted within 180 days of the release of the public verification package, and in any case no later than 30 June 2027, with the deletion recorded in a signed receipt appended to that package.

7.7 Reviewer exposure

Exposure minimisation also determined the reviewer count. A single trained reviewer examined only the material required for validation; adding reviewers would have widened access to explicit names and aliases in metadata and to biometric crops derived from sensitive content. We treat the resulting absence of inter-rater reliability as a methodological limitation, not as evidence of accuracy. It is also a standing position rather than a scheduling gap: we will not put the study material in front of any additional person, including metadata-only review of titles and tags. Reliability work on the production-identifier boundary will therefore be intra-rater (the same reviewer, blind, under fresh opaque identifiers after a time gap), and carries the memory caveat that design implies. Translation and codebook consultation on difficult audit items exposed selected metadata to an AI assistant, and the cross-provider sensitivity pass processed the restricted metadata with `store:false`; no source text is retained in its public aggregate outputs. The same minimisation principle governs release: raw metadata, URLs, uploader identifiers, face crops, embeddings, review-image paths, and reversible linkage keys will not be published. The study deliberately excludes any public output that would identify or expose a person in the sampled content.

7.8 Incident handling

Three defects were found during the study: a parser fault (tnaflix uploaders), contaminated validation rows (spankbang), and an interface ambiguity in the pair review. Each was archived, repaired deterministically, re-adjudicated where decisions were affected, and documented with machine-readable receipts rather than silently fixed (§3.5, Appendix C). We consider disclosure of such corrections part of the governance of sensitive measurement, and the verification package preserves their change trail.

8 Conclusion

We measured two properties of intimate imagery across four adult platforms. The automated analysis finds no searchable person or production identifier in 48.1% of clips pooled across the three uniformly sampled platforms and 56.1% in the observed four-platform sample after including erome’s non-uniform uploader walk. The latter has a 53.0–56.1% self-post sensitivity range. The expanded reviewer-led audit estimates 43.9% on the uniform three and 52.6% overall; GPT-5.6

Sol estimates 40.5% and 49.2%. This convergence supports “about half”, while neither design establishes a general majority. For the ordinary-user upload accounts encountered through that clip sample, the study reports that 90% of the 1,943 accounts crawled to at least ten posts depict ten or more biometrically distinct people, a figure essentially unchanged by inverse-multiplicity reweighting (89.5%). The first result qualifies metadata search; the second qualifies account-by-account browsing. Neither result measures consent, authorship, or all possible discovery routes.

The conclusion is a reframing. The persistent problem in non-consensual intimate imagery is commonly cast as one of takedown, of how quickly platforms act once notified. These measurements show that discovery deserves separate evaluation: metadata often supplies no searchable person or production identifier, and many encountered uploader catalogues are not organised around one depicted person. Improving removal remains necessary, but it cannot by itself solve failures that occur before a report.

Data and Code Availability

The complete study record exists across restricted local and GPU-host storage. It includes the 3,950 per-clip metadata and classification records; both initial 100-item reviewer-audit files; the expanded 400-item audit, its locked private mapping, the completed tnaflx correction review and merge receipts, corrected de-identified collapsed labels, and weighted analysis; the complete corrected GPT-5.6 Sol robustness output and aggregate provenance receipts; 2,497 account-level Finding 2 rows, including 742 clean spankbang re-crawls; the repaired 50-account repeatability record, the 1,654-medoid candidate-pair provenance and review outputs, the archived predecessor, and a de-identified correction receipt; the completed 469-pair de-identified verdict output and aggregate completion receipt; and the calibration summaries. The exact Sonnet and Opus model identifiers were not recorded and cannot be recovered from the available logs.

These records are not all publication-safe. Source metadata can contain explicit names or aliases linked to adult content, and the account and validation files can contain uploader handles, URLs, crop paths, or other linkage material. Raw metadata, direct or reversible identifiers, face crops, embeddings, and review-image paths will therefore not be published. The versioned public verification package published alongside this preprint is limited to analysis code, prompts and codebooks, sampling procedures, aggregate tables, calibration summaries, synthetic fixtures, the expanded audit’s collapsed findable/unfindable/ambiguous labels stripped of source text and linkage keys, safe survey-design fields needed to recompute the weighted audit, safe cross-provider aggregates and provenance, and shuffled verdict-only audit tables, including the de-identified candidate-pair verdict output and aggregate completion receipt. It is available at <https://qayemdigital.com/findable-verification.zip>; its README gives a clean-room verification command and its checksum list hashes every file. Version 1.2 of the archive has SHA-256 7661ddb2a3c3b30768da524123159620052ca797553817b81def84c5e8f53ff2. Code is released under the MIT License and de-identified data and documentation under CC BY 4.0. This paper is licensed under Creative Commons Attribution 4.0 International (CC BY 4.0). Sensitive artifacts follow the retention and deletion schedule in §7. Requests for additional derived data will be considered only under controlled access where disclosure risk can be adequately managed (research@qayemdigital.com).

Authorship and accountability. This study was conducted by QayemDigital Research and is published under organisational authorship. QayemDigital does not publish the names of individual staff. That is a standing privacy practice of the organisation, not a claim about any particular risk to the people who did this work. Accountability is

preserved through the correspondence address, which reaches the responsible researchers, and through the verification package and correction receipts described above.

Competing interest. QayemDigital develops face-search technology. The study’s authors therefore have a commercial interest in the problem of biometric discovery described here. The study does not evaluate or endorse a QayemDigital product. Definitions were locked before analysis, results are checked by a model-blind reviewer-led audit and a recorded exact-model cross-provider pass, and the negative scope of the claims is stated explicitly (§5).

Appendices

A. Threshold calibration. Figure 1 shows the cosine-similarity distributions for same-performer and different-performer face pairs on the real, performer-labelled set (85 performers drawn from the platform’s performer directory; 2,567 same-performer and 237,904 different-performer pairs). The distributions show substantial separation, with median 0.44 (same) versus 0.03 (different), and the operating threshold 0.42 sits in the gap between them. At 0.42, 31 different-performer pairs match (a false-match rate of 1.3×10^{-4} , roughly half the rate at 0.38), all of them similar-looking performers. Because comparisons are clustered within 85 performers, this is an empirical pair-level false-match rate, not an independent-observation confidence estimate. The operating point deliberately favours precision over pairwise recall. The 0.57 genuine-match rate on degraded frame-to-frame pairs does not estimate recall of the account-level ≥ 10 classifier. Transitive multi-frame evidence can recover some borderline matches, but account-level recall is not directly measured; the bounded candidate review probes the resulting over-splitting risk. The performers were labelled automatically by anchoring to each one’s official reference photo; on the two performers a reviewer could personally identify, 105 of 108 automatically labelled faces (97%) were confirmed correct.

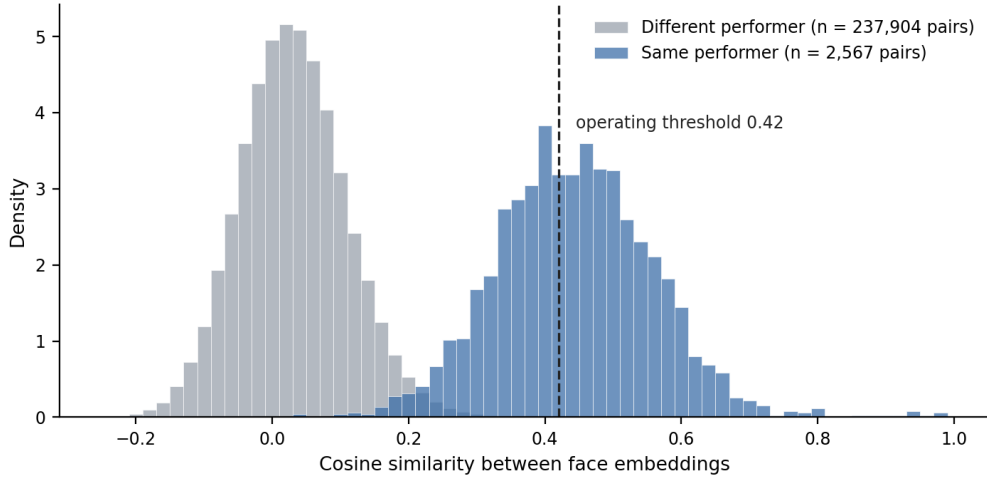


Figure 1: Same-performer versus different-performer cosine similarity on real, performer-labelled adult frames. The operating threshold (0.42) sits in the gap between the two distributions.

B. Classification codebook. The locked searchability definition (§3.3), illustrated with synthetic rather than source identifiers. Searchable by name: metadata naming the depicted person (for example, “Performer Alias: scene title”); a single alias can qualify when the metadata presents it as the depicted person’s name. Searchable by production identifier: a recognisable studio, series, film, scene, or catalogue key (for example, “STUDIO-123”). Unsearchable: metadata that describes

only an act, scenario, body type, date, or generic exact title. Searchability requires a searchable person or production identifier in the metadata, but does not establish that the depicted person knows the term, that it is sufficiently specific, or that a search engine will retrieve the item.

Worked boundary cases (synthetic) record where the codebook is hardest to apply. “Amateur compilation #12” is unsearchable: the number enumerates, it does not key a catalogue. “STUDIO-123 Scene 4” is searchable: the prefix exposes a catalogue family. A bare channel or uploader name is not, by itself, an identifier of the depicted person; it qualifies only when the metadata presents it as the depicted person’s name. A catalogued film code qualifies even without a performer name. The contested-code stratum (Appendix D) shows this boundary is the least stable across instruments; the identified next steps are hardening the codebook with cases like these and a blind intra-rater re-test by the same reviewer (§7).

C. Reproducibility and corrections. The repaired stratified 50-account re-fetch has 15 low-count, 20 mid-count, and 15 high-count accounts (recount $r = 0.997$, median relative difference 0%, and 48/50 within 25%). After the “recommended”-strip defect was found (a pre-clean crawl helper could ingest a page’s “recommended” strip, attributing other pages’ content to an account), all 742 spankbang study accounts were re-crawled with strip exclusion validated by cross-account overlap checks. A later hash audit found that the validation subset still contained 15 pre-clean spankbang rows; those rows and the associated qualitative review were withdrawn, archived, and replaced deterministically from the clean crawl, preserving the 15/20/15 low/mid/high allocation. In the repaired record no exact spankbang medoid hash recurs across accounts. The replacement human review adjudicated all 469 preselected medoid pairs across 25 accounts: 445 different, 11 same, and 13 uncertain. Confirmed edges reduced nine account counts by 11 identities in total (maximum three); no actual or possible ≥ 10 classification changed. The candidate screen is not exhaustive. An early navigation/persistence mismatch was repaired from deterministic image-request logs after the reviewer clarified that *Next* had meant *different*; a one-time migration restored 240 unsaved decisions without overwriting seven explicit saves, a 34-position range (global positions 237–270) was later reset and re-adjudicated, and 230 migrated calls remain in the final output alongside 29 form submissions and 210 post-fix one-click *different* calls. Versioned archives and receipts preserve the transition.

For Finding 1, the tnaflx uploader parser was found to be profile-only and attribute-order dependent, capturing a related-video uploader on many pages. The correction preserved the original 1,000-row sample and produced a restricted per-page ledger, a repaired sample, and an aggregate receipt without handles or URLs. The re-fetch resolved 949 main uploaders, found 49 live pages with no main-uploader badge, and left two current 404s; of the 1,000 pages, 998 supplied a usable corrected state. All 100 tnaflx rows in the locked 400-item audit were usable; 91 had an uploader value changed or cleared, and all 100 were re-adjudicated under newly shuffled opaque IDs while the other 300 decisions were retained. Seventeen reason-level and 11 collapsed decisions changed. The historical Sonnet/Opus pass was not reconstructed because its exact model identifiers were not preserved; the corrected reviewer and recorded exact-model GPT-5.6 Sol passes carry the direct post-repair checks. The correction crosswalk and source metadata remain restricted. The completed correction export, merge receipt, corrected de-identified labels, uploader-status sensitivity, and the exact-model tnaflx-only GPT-5.6 Sol rerun preserve the change trail.

D. Expanded audit detail. The 400-item audit was locked before review with 150 named controls, 50 confirmed-code records, 50 contested-code records, and 150 records with no historical identifier. Table 5 reports each stratum’s population share, reviewer-led unsearchable estimate, and

raw agreement. Nine items were abstentions. Exact platform-by-stratum inclusion weights, rather than the rounded shares shown here, produce the prevalence and weighted agreement in §4.3. The interface’s reason categories were not consistently applied in the first 100 decisions, so all name, production, and both responses are collapsed to findable; no reason-specific prevalence is reported.

Table 5: Expanded searchability audit by historical-pipeline stratum. “Reviewer unsearchable” and “Agreement” use decided items; nine of 400 items were abstentions. The final row gives raw and survey-weighted agreement.

Stratum	Audited	Population share	Reviewer unsearchable	Agreement
Named (control)	150	40.5%	5.4%	94.6%
Confirmed code	50	3.4%	12.2%	88.0%
Contested code	50	2.1%	39.6%	39.6%
No identifier	150	54.0%	91.0%	91.1%
Overall / weighted	400	100.0%	52.6%	85.7% / 91.3%

References

- [1] TAKE IT DOWN Act, Pub. L. No. 119-12 (2025), United States.
- [2] Online Safety Act 2023, c. 50, United Kingdom.
- [3] Regulation (EU) 2022/2065 of the European Parliament and of the Council (Digital Services Act), 2022.
- [4] C. McGlynn and E. Rackley, “Image-Based Sexual Abuse,” *Oxford Journal of Legal Studies*, vol. 37, no. 3, 2017.
- [5] D. K. Citron and M. A. Franks, “Criminalizing Revenge Porn,” *Wake Forest Law Review*, vol. 49, 2014.
- [6] N. Henry and A. Powell, “Technology-Facilitated Sexual Violence: A Literature Review of Empirical Research,” *Trauma, Violence, & Abuse*, vol. 19, no. 2, 2018.
- [7] Y. Ruvalcaba and A. A. Eaton, “Nonconsensual Pornography Among U.S. Adults: A Sexual Scripts Framework on Victimization, Perpetration, and Health Correlates for Women and Men,” *Psychology of Violence*, vol. 10, no. 1, 2020.
- [8] J. Guo, J. Deng, A. Lattas, and S. Zafeiriou, “Sample and Computation Redistribution for Efficient Face Detection,” *International Conference on Learning Representations (ICLR)*, 2022.
- [9] J. Deng, J. Guo, N. Xue, and S. Zafeiriou, “ArcFace: Additive Angular Margin Loss for Deep Face Recognition,” *IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2019.
- [10] X. An et al., “Partial FC: Training 10 Million Identities on a Single Machine,” *IEEE/CVF International Conference on Computer Vision Workshops (ICCVW)*, 2021.
- [11] PimEyes, “How to use PimEyes: a brief guide,” pimeyes.com/blog. Accessed July 2026. pimeyes.com/en/blog/how-to-use-pimeyes-a-brief-guide.
- [12] UK Information Commissioner’s Office, “How do we process biometric data lawfully?” *Biometric recognition guidance*. Accessed July 2026. ico.org.uk/biometric-recognition-guidance.
- [13] J. M. Urban, J. Karaganis, and B. L. Schofield, “Notice and Takedown in Everyday Practice,” UC Berkeley Public Law Research Paper No. 2755628, 2017.
- [14] T. Gillespie, *Custodians of the Internet: Platforms, Content Moderation, and the Hidden Decisions That Shape Social Media*. Yale University Press, 2018.

- [15] S. T. Roberts, *Behind the Screen: Content Moderation in the Shadows of Social Media*. Yale University Press, 2019.
- [16] D. K. Citron, “Sexual Privacy,” *Yale Law Journal*, vol. 128, no. 7, 2019.
- [17] N. Henry, C. McGlynn, A. Flynn, K. Johnson, A. Powell, and A. J. Scott, *Image-Based Sexual Abuse: A Study on the Causes and Consequences of Non-Consensual Nude or Sexual Imagery*. Routledge, 2020.
- [18] Microsoft, “PhotoDNA.” Accessed July 2026. microsoft.com/en-us/photodna.
- [19] SWGfL, “StopNCII: Stop Non-Consensual Intimate Image Abuse.” Accessed July 2026. stopncii.org.
- [20] National Center for Missing & Exploited Children, “Take It Down.” Accessed July 2026. [takeit-down.ncmec.org](https://takeitdown.ncmec.org).
- [21] K. Hill, *Your Face Belongs to Us: A Secretive Startup’s Quest to End Privacy as We Know It*. Random House, 2023.
- [22] H. Ajder, G. Patrini, F. Cavalli, and L. Cullen, “The State of Deepfakes: Landscape, Threats, and Impact,” Deeptrace Labs, 2019.
- [23] P. Vallina, Á. Feal, J. Gamba, N. Vallina-Rodriguez, and A. Fernández Anta, “Tales from the Porn: A Comprehensive Privacy Analysis of the Web Porn Ecosystem,” *ACM Internet Measurement Conference (IMC)*, 2019.
- [24] D. Dittrich and E. Kenneally, “The Menlo Report: Ethical Principles Guiding Information and Communication Technology Research,” U.S. Department of Homeland Security, 2012.