

Contents

- Fault Amplification in Networked ALPR Systems
 - Data Error, Automated Suspicion, and the Governance of Flock Safety
 - Abstract
 - 1. Introduction
 - 2. Research Questions
 - 3. Prior Research
 - 3.1 Targeted value, mixed general effects
 - 3.2 Error is sociotechnical
 - 3.3 Automation bias and coercive escalation
 - 3.4 Aggregated movement and the Fourth Amendment
 - 4. Data and Methods
 - 4.1 Source hierarchy
 - 4.2 Deployment and volume datasets
 - 4.3 Documented-case registry
 - 4.4 Deterministic collision experiment
 - 4.5 Synthetic OCR experiment
 - 4.6 Why FBI UCR is excluded
 - 5. Results
 - 5.1 A large but differently measured footprint
 - 5.2 Mass capture and sparse contemporaneous hits
 - 5.3 Deterministic fault amplification
 - 5.4 Open-model degradation results
 - 5.5 Documented pathways of harm
 - 5.6 Utility does not answer governance
 - 6. Legal and Institutional Analysis
 - 6.1 Collection, query, and seizure are distinct events
 - 6.2 The local-control fiction
 - 6.3 Auditability is necessary but insufficient
 - 7. Policy Framework: Twelve Breakpoints
 - 8. Limitations
 - 9. Conclusion
 - Reproducibility Appendix
 - A. Commands
 - B. Primary generated outputs
 - C. Interpretation rules
 - D. Research ethics

Fault Amplification in Networked ALPR Systems

Data Error, Automated Suspicion, and the Governance of Flock Safety

Flockwatch Research Project · July 2026

Abstract

Automated license-plate readers (ALPRs) are frequently evaluated as sensors: a camera either recognizes a plate or it does not. That framing is incomplete. Modern commercial systems operate as networked decision infrastructures in which optical recognition, police hot lists, cross-jurisdictional sharing, discretionary database queries, and high-risk vehicle-stop tactics interact. A small defect at any layer may therefore propagate farther and produce more serious consequences than the original error would suggest. This study calls that process **fault amplification**.

Using only public sources, we combine four forms of evidence: a national deployment inventory from the Electronic Frontier Foundation's Atlas of Surveillance; openly licensed camera-mapping and agency-portal datasets published by Eyes Off Indiana; California public-records datasets containing billions of ALPR detections and recorded hot-list hits; and an evidence-graded registry of publicly documented false stops, insider misuse, interstate sharing, reproductive-health searches, immigration searches, and protest-related queries. We also conduct a privacy-safe local experiment comprising a deterministic plate-key collision model and 600 inferences by an open-source OCR model on synthetic U.S.-style plate crops.

The public data show scale without establishing a national harm rate. The Atlas contains 4,084 records classified as ALPR deployments, 2,629 of which name Flock Safety; these are agency or technology records, not camera counts. A separate OpenStreetMap-derived census reports 113,963 publicly mapped cameras. California records contain 3.224 billion detections in 2021–2022 and 3.557 million recorded hits, a 0.110% hit share. The inverse should not be equated with innocence, but it demonstrates that the infrastructure primarily records vehicles not producing contemporaneous hot-list hits. In the local laboratory, deleting a two-digit middle token caused 100 distinct synthetic identifiers to collapse to one match key. Under blurred quarter-resolution rendering, the open OCR model preserved that middle token in one of 100 trials. These synthetic results demonstrate a mechanism, not Flock's proprietary error rate.

The findings support a governance model that treats ALPR output as an investigative lead rather than self-authenticating suspicion. Effective controls must interrupt the entire chain: source-data validation, full-identifier matching, independent officer verification, short retention, warrant requirements for historical searches, restricted sharing, tamper-evident audits, meaningful remedies, and independent effectiveness evaluation. ALPRs may provide targeted operational value, particularly in stolen-vehicle investigations. That utility does not eliminate the need to govern how uncertain machine observations become police force.

Keywords: automated license-plate readers; Flock Safety; surveillance; automation bias; false positives; police databases; Fourth Amendment; public records; fault amplification

1. Introduction

On a June weekend in 2026, automotive journalist Joel Feder and his wife were boxed in by four police cruisers in Plymouth, Minnesota. Officers had been following an automated stolen-vehicle alert. Feder's manufacturer plate included the sequence 34 10 DTM ; the lost plate entered by another agency was 34 03 DTM . According to Feder's first-person account, the Los Angeles Police Department's record omitted one middle number group while the Flock camera's read omitted the other. Both became 34 DTM , and the shortened value propagated through a national alert network (Feder 2026).

The event was not a single bad read. It was a chain:

1. a source record lost information;
2. a sensor output lost different information;
3. the two lossy strings were treated as equivalent;
4. a network distributed the alert beyond the originating jurisdiction;
5. officers interpreted the alert as sufficiently reliable to organize a high-risk stop; and
6. the driver bore the physical and psychological risk of resolving the system's uncertainty.

That sequence motivates this paper's central claim: **the danger of networked surveillance is not only that it observes at scale, but that it amplifies faults across technical, organizational, and coercive layers.** The camera is one component. The operative system includes recognition software, vendor normalization, state and federal hot lists, agency configuration, national sharing, query interfaces, audit practices, and officer decision-making.

This framing matters because public debate often collapses into a binary. Supporters point to recovered vehicles, missing-person cases, and investigative leads. Critics point to mass collection and chilling effects. Both can be true. The empirical question is not whether ALPRs ever help police; peer-reviewed research indicates targeted operational value under some conditions. The governance question is whether an investigative benefit justifies collecting and networking the

movements of millions of uninvolved drivers while allowing errors and discretionary queries to trigger coercive action.

We make five contributions. First, we distinguish deployment records, mapped cameras, detections, hits, searches, stops, arrests, and convictions rather than treating them as interchangeable evidence of effectiveness. Second, we construct a public-source fault taxonomy spanning capture, data entry, matching, sharing, human verification, and enforcement response. Third, we reproduce the Feder-style information-loss mechanism in a deterministic synthetic registry. Fourth, we pair nationally scoped public data with a documented-case registry while refusing to infer prevalence from selected incidents. Fifth, we derive a twelve-control governance framework aimed at breaking fault chains before uncertain data become force.

2. Research Questions

The study addresses four questions:

RQ1 — Scale. What does available public data establish about the documented footprint and volume of U.S. ALPR infrastructure?

RQ2 — Fault amplification. Through what mechanisms can recognition errors, partial identifiers, stale hot lists, permissive sharing, or insider queries produce consequences beyond the original defect?

RQ3 — Evidence of utility and harm. What do independent evaluations, public records, court opinions, audits, and documented incidents support—and what do they not support—about effectiveness and risk?

RQ4 — Governance. Which controls can interrupt the path from uncertain observation to historical tracking or high-risk police intervention?

We do not estimate a national false-positive rate. No public, representative corpus links modern commercial ALPR reads to verified ground truth, downstream stops, and dispositions across jurisdictions. We also do not estimate Flock Safety's proprietary OCR accuracy. The local experiment measures an open model on synthetic stimuli designed to isolate a public failure mechanism.

3. Prior Research

3.1 Targeted value, mixed general effects

Independent research does not support the claim that ALPRs are categorically useless. Taylor, Koper, and Woods found that an ALPR deployment increased scanning, stolen-vehicle hits, arrests, and recoveries in an Arizona experiment (Taylor et al. 2012). Koper, Taylor, and Woods

found conditional residual crime effects under particular directed-patrol strategies (Koper et al. 2013). A later study of a large fixed network found potential gains in investigative clearance for some offenses while emphasizing that use was concentrated in a small share of cases and that detective discretion prevented clean causal attribution (Koper and Lum 2019).

The strongest evidence is narrower than vendor rhetoric. A randomized hot-spot experiment by Lum and colleagues found no general or auto-crime deterrent effect under the tested low-dosage deployment (Lum et al. 2011). Implementation research further shows that technology does not automatically change organizational practice or outcomes (Lum, Koper, and Willis 2017). Flock Safety's extrapolations about crimes "solved" derive from customer surveys and corporate analysis, not independent causal evaluation. We therefore treat recovery and investigation as plausible targeted benefits while rejecting raw alert, query, or customer-testimonial counts as proof of population-level crime reduction.

3.2 Error is sociotechnical

ALPR error is often presented as character recognition error. The legal record demonstrates a broader chain. In *Green v. City and County of San Francisco*, an ALPR allegedly read 5SOW350 as 5SOW750, officers did not visually verify the plate, and Denise Green was subjected to a high-risk stop. The Ninth Circuit held that a jury could find the seizure unreasonable given known system fallibility and the failure to corroborate the alert (751 F.3d 1039).

Other publicly reported failures include stale stolen-vehicle records, wrong-state and wrong-vehicle matches, and character substitutions. In Aurora, Colorado, Brittney Gilliam and children were detained at gunpoint after a plate associated with a stolen motorcycle was matched to an SUV from a different state; the city later reached a \$1.9 million settlement (The Guardian 2024). In New Mexico, reporting described a 2 / 7 confusion followed by a gunpoint stop with children present (KOAT 2022). These incidents differ in vendor, software, and legal posture. They are not a rate study. Together with *Green*, they demonstrate recurring failure classes.

3.3 Automation bias and coercive escalation

Automation-bias research describes commission errors in which people follow incorrect machine advice and omission errors in which they fail to act because an automated system does not alert. The evidence originated in aviation and clinical decision support rather than ALPR field experiments (Goddard, Roudsari, and Wyatt 2012). Its relevance here is mechanistic, not a claim that officers have been experimentally shown to behave identically to pilots or clinicians.

An ALPR alert compresses uncertainty into a simple operational cue: hit/no hit, often paired with a high-risk category such as stolen vehicle or wanted person. Under time pressure, that cue may replace independent information search. Professional guidance therefore requires visual and database verification before a stop when circumstances permit. Fault amplification occurs when the system supplies a strong action signal while obscuring how weak or lossy the underlying match is.

3.4 Aggregated movement and the Fourth Amendment

A single observation of a visible plate on a public road has generally received limited constitutional protection. Network aggregation changes the practical question. *United States v. Jones* and *Carpenter v. United States* recognize that prolonged technology-assisted location tracking can reveal the whole of a person's movements even when individual observations occur in public (*Jones*, 565 U.S. 400; *Carpenter*, 585 U.S. 296).

In *Commonwealth v. McCarthy*, Massachusetts's highest court held that the limited ALPR use in the record was not a search while expressly recognizing that sufficiently widespread surveillance could implicate constitutional privacy (142 N.E.3d 1090). The Ninth Circuit rejected suppression in *United States v. Yang* on the facts before it (958 F.3d 851). No Supreme Court decision squarely resolves warrantless historical searches of dense, privately operated national ALPR networks. The current landscape is therefore a patchwork of narrow observation doctrine, aggregation principles, state statutes, vendor contracts, and agency policy.

4. Data and Methods

4.1 Source hierarchy

We used the following hierarchy: court opinions, statutes, official audits, and government datasets; peer-reviewed studies; agency-released public records; university public-records research; first-person accounts with underlying documents; reputable investigative journalism; advocacy analysis linked to records; and vendor material. Source class is retained because corporate, advocacy, and police sources have different incentives.

All reproducible inputs are enumerated in `sources/data_sources.json`; downloaded content is hashed in `data/provenance/manifest.json`. Original third-party datasets are fetched rather than silently republished when their redistribution terms are unclear.

4.2 Deployment and volume datasets

The EFF Atlas of Surveillance snapshot contained 15,071 total technology records. We selected rows whose technology field was `Automated License Plate Readers`, producing 4,084 records. An Atlas row documents an agency or jurisdiction's use of a technology based on a public source. It is **not** a camera. We classified a row as Flock when the vendor field contained `Flock Safety`.

For camera counts we used Eyes Off Indiana's national state rankings, derived from OpenStreetMap camera observations with 2024 Census population denominators. Because mapping depends on volunteer and public-record coverage, the result is a documented lower-bound census, not a complete vendor inventory. We used its monthly Indiana series to describe public documentation growth, not to assert exact installation dates.

For activity volume we used EFF's California ALPR detections-and-hits file for 2021–2022 and the reliable-agency worksheet from *Data Driven 2* for 2018–2019. We converted spreadsheet-formatted counts to integers and summed detections and recorded hot-list hits. A “hit” is a database event, not necessarily a correct match, stop, arrest, charge, conviction, recovered vehicle, or prevented offense. A “non-hit” is not necessarily irrelevant to later historical investigation.

Eyes Off Indiana's portal dataset contains self-reported 30-day detection volumes, camera counts, and agency names visible in public Flock transparency portals. These values are useful for order-of-magnitude exposure, but agencies may update them on different schedules and the sample is not national.

4.3 Documented-case registry

We constructed an evidence-graded registry of 13 public incidents and governance failures. Grade A requires a primary or official source plus corroboration; Grade B indicates strong multi-source journalism or public-records analysis. Pending complaints and criminal charges are labeled as allegations rather than adjudicated findings. The registry is purposive: incidents enter because they are documented and relevant to a failure class. It cannot estimate frequency, relative risk, demographic disparity, or vendor market performance.

4.4 Deterministic collision experiment

The first laboratory experiment isolates information loss. We generated 100 synthetic identifiers from `34 00 DTM` through `34 99 DTM`. The tested transformation removes the two-digit middle token and non-alphanumeric separators. Every full identifier therefore maps to `34DTM`.

This is deliberately transparent rather than probabilistic. It asks: if a system removes the distinguishing token in this format, how many candidate identifiers become indistinguishable? The answer follows from the defined registry and is protected by unit tests, including the motivating `34 03 DTM` / `34 10 DTM` pair.

4.5 Synthetic OCR experiment

We rendered 100 privacy-safe synthetic plate crops using Pillow and the open Noto Sans Mono font. Each image contained large `34` and `DTM` tokens with a smaller two-digit middle token. We generated six conditions: native, half resolution, quarter resolution, half resolution plus blur, quarter resolution plus blur, and low-contrast quarter resolution. We then ran Fast Plate OCR 1.1.0's MIT-licensed `cct-xs-v2-global-model` locally on CPU, for 600 total inferences.

We report exact full-string recovery and whether the middle token appeared in the model output. This stimulus is intentionally difficult and out of distribution: it resembles a stacked or hierarchical identifier rather than a conventional single-line plate. Consequently, native performance is not a general baseline. The experiment tests whether the distinguishing small

token can disappear under controlled degradation; it does not estimate the accuracy of Flock Safety, Fast Plate OCR on ordinary plates, or operational ALPR systems.

4.6 Why FBI UCR is excluded

The Uniform Crime Reporting program measures offenses known to police and related arrest or clearance statistics. It does not record how many plates a vendor scans, how long observations are retained, how often a hot list is stale, whether an alert was visually verified, whether a database query concerned immigration or protest activity, or whether an innocent driver was stopped because two identifiers collapsed into one. UCR may be useful in a separately designed crime-outcome evaluation with credible treatment and comparison groups. It cannot serve as a direct measure of ALPR exposure or harm, so it is not used here.

5. Results

5.1 A large but differently measured footprint

The Atlas contained **4,084 documented ALPR records** across the fifty states and the District of Columbia. Flock Safety appeared in **2,629 records**, or **64.4%**. This indicates strong representation in EFF's documented adoption corpus. It is not a market-share estimate: records vary in age and granularity, one agency may operate many cameras, and unknown-vendor records remain.

The separate public map reported **113,963 cameras** across the fifty-state rankings file, equivalent to **33.6 mapped cameras per 100,000 residents** using its population denominators. The two datasets answer different questions. Atlas is a sourced agency-technology inventory. The map is a camera-location census. Their counts should never be added.

Indiana illustrates the pace at which public documentation can expand. Its monthly series increased from one mapped camera in December 2022 to **3,170** in July 2026. This is a 3,169-record increase in mapped cameras, not a verified installation cohort and not evidence that all were installed during the interval.

5.2 Mass capture and sparse contemporaneous hits

Across the California 2021–2022 file, agencies reported **3,224,074,784 detections** and **3,556,754 hits**, a hit share of **0.110%**. The corresponding 99.890% non-hit share should not be labeled “innocent drivers”: a non-hit observation can later support a historical investigation, and some drivers or vehicles may be associated with offenses not represented on a contemporaneous hot list. The ratio nevertheless establishes a structural fact: a hot-list system records orders of magnitude more observations than it flags at capture time.

In the reliable-agency worksheet for 2018–2019, 63 agencies reported **1,681,232,745 detections** and **873,788 hits**, a hit share of **0.052%**. Differences between periods may reflect agency

composition, data completeness, deployment growth, hot-list policy, or counting practice; they are not a causal time trend.

The Indiana portal sample contained 18 reporting agencies, 189 cameras, and **3,920,631 detections over 30 days**, or approximately **691.5 detections per camera per day**. It is a convenience sample of public portals, not a national average. Its importance is interpretive: even modest local networks can create millions of records per month.

5.3 Deterministic fault amplification

The collision experiment produced the intended 100:1 compression. All 100 full identifiers mapped to one shortened key. That does not show a 100-person false-match event occurred in Feder's case. It shows what a database loses when a two-digit distinguishing token is removed: the query no longer names one member of the synthetic family.

The distinction between error and ambiguity is crucial. A substitution such as **2** for **7** may point to the wrong single candidate. Deletion can create a non-unique query that should be represented as an ambiguity set. If an interface instead presents one alert without showing the lost characters or candidate count, it converts missing information into false confidence.

5.4 Open-model degradation results

The open OCR model exactly recovered 20 of 100 native synthetic identifiers and preserved the middle token in 24. At half resolution, exact recovery rose to 30 and token preservation to 31, likely because resizing changed antialiasing in a way that better matched the model's training distribution. This non-monotonicity reinforces that the experiment is not an operational accuracy benchmark.

At quarter resolution, exact recovery fell to 3% and token preservation to 10%. With half-resolution blur, token preservation fell to 2%. Under blurred quarter resolution and low-contrast quarter resolution, the model preserved the middle token in **one of 100** images and exactly recovered none.

The finding is bounded: in this synthetic layout, an open OCR model frequently dropped the small token under degradation. Combined with the deterministic experiment, it demonstrates how token loss can enlarge a match set. It does not establish that Flock uses the same model, preprocessing, imagery, normalization, or frame aggregation.

5.5 Documented pathways of harm

The registry organizes incidents by mechanism rather than tallying them as a rate.

Sensor and match error. *Green*, Gilliam, Gonzales, Feder, and the pending Upchurch complaint describe wrong-character, wrong-state, wrong-vehicle, stale-record, or partial-identifier failures

followed by police intervention. The recurring control failure is not merely bad OCR; it is failure to resolve readily available contradictions before escalation.

Insider misuse. Georgia investigators have brought cases alleging personal stalking or harassment through ALPR access. Audit logs enable later detection, but detection after repeated searches is not prevention. Charges remain allegations until adjudicated.

Cross-jurisdictional purpose drift. Public-records research in Colorado, Illinois, and Washington documented immigration-related searches or federal access through local networks. A Texas abortion-related investigation generated a nationwide search across jurisdictions with conflicting laws. These cases challenge the claim that each locality controls only its own cameras: network search allows one agency's policy choice to expose another community's data.

Protected activity and discriminatory queries. EFF's analysis of released logs identified protest-related and Romani-related searches. The records document queries; they do not prove how every result was used or establish a population-level chilling effect. They do show that free-text purpose fields and broad national access permit uses far beyond stolen-car recovery.

5.6 Utility does not answer governance

A useful stolen-vehicle alert and an unlawful personal search can occur in the same system. Aggregate success counts cannot answer whether historical location searches need warrants, whether non-hit observations should be retained, whether one state may search another's cameras for conduct lawful there, or whether a gunpoint stop was reasonable after an obvious mismatch.

This separation also clarifies evaluation. Camera vendors often report investigations “supported” or crimes “solved,” but those categories may count duplicate leads, cases where other evidence was decisive, or customer attribution without a comparison group. Independent studies support a narrower proposition: targeted deployments can increase operational outputs and may improve some investigations. Evidence for broad deterrence is mixed. Governance must therefore be justified by measured outcomes and constrained regardless of usefulness.

6. Legal and Institutional Analysis

6.1 Collection, query, and seizure are distinct events

Constitutional analysis should separate three acts. First is collection of a visible plate at a public location. Second is aggregation and query of historical observations. Third is the stop, search, or use of force triggered by an alert. A court may find no search in a sparse historical query yet still find a resulting seizure unreasonable when officers ignore known system fallibility, as the contrast between *McCarthy* or *Yang* and *Green* demonstrates.

Networked commercial storage complicates the public-observation analogy. A person expects individual drivers or officers to see a vehicle on a road. That does not answer whether thousands of fixed sensors may create a searchable longitudinal dossier. *Carpenter* teaches that aggregation can be constitutionally different from its individual components, but current doctrine has not produced a uniform ALPR warrant rule.

6.2 The local-control fiction

Flock Safety emphasizes customer control over sharing. Public audits reveal several routes around that simple model: direct bilateral sharing, misconfigured broad access, federal pilots, nationwide queries by other local officers, and informal searches performed on behalf of an outside agency. The University of Washington Center for Human Rights describes front-door, back-door, and side-door access ([UWCHR 2025](#)).

The distinction matters because a city can adopt a restrictive policy while its data remain discoverable through another agency's query. Conversely, an officer can search a national network even when the officer's own jurisdiction owns few cameras. Governance must attach to data wherever it travels, not only to the agency that purchased the hardware.

6.3 Auditability is necessary but insufficient

Audit logs have enabled many of the strongest public findings in this study. They can identify who searched, when, and with what written reason. Yet a free-text reason such as “missing person” may conceal the actual investigative purpose; users may know which language avoids scrutiny; and audits may occur months after harm. Logs also do not prevent a false alert from producing a stop in seconds.

Audit design must therefore be paired with access restriction, immutable purpose codes, case linkage, anomaly detection, supervisory review, public aggregate reporting, and consequences. A system that records abuse but rarely checks the record is observable, not accountable.

7. Policy Framework: Twelve Breakpoints

The objective is not to make every sensor perfect. It is to ensure that no single imperfect component can authorize serious action by itself.

1. **Independent verification.** Before a stop, verify the full plate, issuing state, plate type, vehicle make/model, and current hot-list status.
2. **High-risk-stop threshold.** Treat an ALPR alert as an investigative lead, not sufficient grounds for a felony stop without corroborating facts.
3. **Source-data integrity.** Preserve full identifiers, validate unusual plate formats, and promptly remove recovered vehicles and canceled alerts.

4. **Uncertainty-preserving interfaces.** Display omitted characters, confidence, candidate-set size, and conflicting vehicle metadata rather than collapsing them into a binary hit.
5. **Short non-hit retention.** Delete observations not linked to a documented investigation on a short statutory schedule.
6. **Warrants for historical movement searches.** Require probable cause for longitudinal or pattern-of-life queries, subject to narrowly defined emergencies.
7. **Purpose limitation.** Prohibit or strictly regulate immigration, reproductive-health, protest, religion, and generalized identity searches.
8. **Network sharing control.** Default to no national sharing; require named agreements, legal review, expiration, and downstream-use restrictions.
9. **Tamper-evident audit.** Record user, case, purpose, scope, results viewed, export, and downstream action; prevent retroactive editing of reasons.
10. **Independent audit and outcome reporting.** Sample searches and stops, investigate outliers, publish correction and misuse statistics, and connect alerts to dispositions.
11. **Meaningful remedy.** Provide notice, correction, complaint, preservation, exclusion, civil-remedy, and discipline pathways.
12. **Transparent procurement and evaluation.** Disclose contracts, retention, sharing, training, error handling, and renewal decisions; pre-register outcome measures and use credible comparison designs.

These controls are complementary. Short retention does not cure a dangerous real-time stop. Officer verification does not cure mass historical tracking. Audit logs do not cure unlawful sharing unless independent institutions inspect them and can impose consequences.

8. Limitations

This study has substantial limitations.

Public-data selection. Atlas and OpenStreetMap coverage reflect what researchers, reporters, agencies, and volunteers can document. Missingness is not random. Camera counts are lower bounds, and vendor representation may be biased toward visible procurements.

Administrative inconsistency. Agencies may define detections, hits, duplicate reads, and reporting periods differently. Hit ratios are descriptive, not comparable measures of accuracy or effectiveness without stronger metadata.

No linked outcomes. The volume datasets do not connect each alert to verification, stop, force, arrest, recovery, charge, dismissal, or correction. Consequently, we cannot calculate positive predictive value at the point of enforcement or an individual risk of erroneous detention.

Incident ascertainment. Public incidents are more likely to be documented when video exists, litigation is filed, advocates investigate, or a reporter is involved. The registry cannot show prevalence or demographic distribution.

Synthetic external validity. The OCR lab uses one open model, one font, one controlled renderer, and a deliberately unusual small-token layout. Commercial systems may use multiple frames, different optics, proprietary models, plate-format priors, and human review. The experiment demonstrates a possible information-loss mechanism only.

Law is changing. Statutes, contracts, vendor features, and litigation evolve quickly. The legal discussion is a July 2026 snapshot, not legal advice.

Causal crime effects. We do not estimate how Flock adoption changes crime. Such a study would require validated deployment timing, outcome definitions, comparison jurisdictions, spillover analysis, and safeguards against simultaneous policy changes. UCR alone would not satisfy that design.

9. Conclusion

Networked ALPR systems are not merely cameras. They are pipelines that transform partial observations into searchable histories and action signals. The Feder incident is analytically important because two different losses of information—one in a police record and one in a camera read—converged on the same shortened key. A national network then gave that weak equivalence operational reach.

Public data show infrastructure at substantial scale and mass capture far beyond contemporaneous hot-list hits. Public incidents and audits document false stops, stalking allegations, purpose drift, and cross-jurisdictional searches. Independent evaluations also show that ALPRs can produce targeted operational benefits. The honest conclusion is therefore neither “the technology never works” nor “successful recoveries settle the debate.” It is that useful systems can still be dangerously governed.

Fault tolerance in other high-stakes domains relies on independent checks, uncertainty disclosure, limited privileges, auditable actions, and fail-safe defaults. Policing should demand no less. When a machine observation can lead to guns, detention, or a retrospective map of a person's movements, uncertainty must slow the system down rather than disappear inside it.

Reproducibility Appendix

A. Commands



B. Primary generated outputs

- `data/derived/results.json` — top-line descriptive results
- `data/derived/atlas_state_records.csv` — Atlas ALPR record counts by state
- `data/derived/atlas_vendor_records.csv` — Atlas vendor-field counts
- `data/derived/state_camera_rankings.csv` — public mapping rankings
- `data/derived/ocr_degradation_results.csv` — all 600 synthetic inferences
- `data/derived/ocr_summary.json` — condition-level OCR summaries
- `data/provenance/manifest.json` — retrieval metadata and SHA-256 hashes
- `sources/incidents.csv` — evidence-graded documented-case registry
- `sources/legal_cases.csv` — case-law ledger
- `sources/policy_controls.csv` — governance framework and sources

C. Interpretation rules

- Atlas record ≠ camera.
- Mapped camera ≠ complete installed inventory.
- Detection ≠ unique vehicle or person.
- Hit ≠ correct match, stop, arrest, recovery, charge, conviction, or prevented crime.
- Non-hit ≠ proof of innocence or irrelevance.
- Documented incident ≠ prevalence.
- Synthetic OCR result ≠ Flock accuracy.

D. Research ethics

The project does not query live ALPR systems, collect private movement histories, publish real plate-owner records, or attempt to identify drivers. Synthetic plate crops are intentionally

fictional. Public audit-log sources are analyzed at the aggregate or already-published incident level.