

Trust in Communications

The Complete Series

*AI can counterfeit the individual.
It cannot counterfeit the relationship.*

Introduction · Seven Papers · Epilogue · Two Companion Volumes

A +Trusted Infrastructure series from ICA AI

Series editor: John Perkins, Chief Executive Officer

Monograph edition · July 2026

Contents & Reading Order

Introduction	What the series argues, and how to read it	3
Paper 1	What Trust Actually Is	7
Companion A	Voice Fraud: The Evidence	12
Paper 2	The Role of Caller Authentication	22
Paper 3	The One Layer That Can Stand Alone.....	26
Paper 4	The Permanent Coverage Gap	33
Paper 5	Even If the Proposed FCC Rules Work Perfectly	39
Companion B	The Proposed FCC Rules	44
Paper 6	Why a Relationship Cannot Be Faked at Scale.....	52
Paper 7	The Compromised Account.....	57
Epilogue	What Changed, What Remains True	61

Introduction

What the series argues, and how to read it

SERIES INTRODUCTION · ICA AI, Inc. · +Trusted Infrastructure · July 2026

A +Trusted Infrastructure series from ICA AI

Series editor: John Perkins, Chief Executive Officer

Communication has a trust problem, and nowhere is it sharper than in voice — where the industry’s fix, proving who is calling, answers a narrower question than the one that matters. Identity can be faked. Relationships can’t. This series asks what establishing trust actually requires, treats caller authentication as a useful but partial tool, and follows the question across seven short papers and two companion volumes. This introduction explains the purpose, the through-line, and how to read them.

Why this series exists

Trust is failing in voice faster than almost anywhere else in communications. Fraud losses keep climbing, unwanted and spoofed calls remain routine, and the tools meant to restore confidence have not closed the gap. The industry’s response has centered on identity — attestation, know-your-customer rules, caller authentication — all aimed at verifying the identity behind a call. That work is real, and this series does not dismiss it. The most heavily mandated of those tools shows the gap is not closing on its own: after five years of mandate, attestation coverage has yet to reach half of all calls — 48.8% as of June 2026 — and the calls it does sign include, at the highest grade, the very traffic it was meant to screen out. The Evidence companion sets out the figures. But it answers the question who is this? when the question people actually need answered is should I trust this? — a question that turns on the relationship behind a call, not only the identity behind it. Those are not the same question, and the distance between them is the subject of everything that follows. Voice is where the series begins, because that is where the failure is sharpest and most current, and where the evidence is public, dated, and testable — which makes it the proving ground for the argument, not its boundary. The same distinction between identity and relationship reaches text, email, and machine-to-machine communication alike.

Why it takes seven papers

The question — is verifying identity enough to establish trust? — has several parts, and answering it means taking them one at a time rather than asserting a conclusion and defending it. Each paper isolates one part. The early papers establish what trust and identity actually are, and where they diverge. The middle papers test whether identity is the thing a trust decision rests on, and whether hardening it — through better attestation or further regulation — can ever be enough. The later papers examine the proposed FCC rules on their best possible day, the economics that keep a relationship from being faked at scale, and the one place where the relationship approach is tested hardest. Read in order, they form a single argument. Read alone, each still stands.

The through-line

One sentence carries the series: caller authentication is useful, but not necessary — and not sufficient — for trust. Useful, because knowing who is calling has real value. Not necessary, because this series argues that trust can be established from the relationship itself — its history, context, and reciprocity — with or without a credential. Not sufficient, because a valid credential on a compromised account passes every identity check and still should not be trusted. Hold those three together and a different picture of the problem emerges: identity is one input to a trust decision, not the decision itself. The purpose of this series is not to ask readers to accept this conclusion, but to examine whether it follows from the evidence presented in the papers that follow.

The seven papers

1. **What Trust Actually Is** — trust is relational and historical; a verified identity is a narrow, recent substitute for it. This paper defines the terms — a relationship as the accumulated record of interaction both parties hold, historical, contextual, and reciprocal — and the rest of the series builds on that definition.
 2. **The Role of Caller Authentication in Establishing Trust** — the genuine, bounded role caller authentication plays — and why “useful” is not “required.”
 3. **The One Layer That Can Stand Alone** — what a trust decision actually rests on once the credential is removed.
 4. **The Permanent Coverage Gap** — why hardening authentication cannot reach the under-resourced edges where fraud routes.
 5. **Even If the Proposed FCC Rules Work Perfectly, Will They Be Enough?** — why compliance, even flawless, cannot be the sole arbiter of trust.
 6. **Why a Relationship Cannot Be Faked at Scale** — the structure and economics that limit impersonation.
 7. **The Compromised Account** — the honest case — what happens when a real account is taken over and the credential is genuine.
- ◆ **Epilogue — What Changed, What Remains True — what the series changed, what it leaves standing, where it stops, and the questions it hands to the field. Read last.**

The two companions

Two companion volumes sit alongside the papers, each meant to be read where it does the most work. Voice Fraud: The Evidence gathers the dated, third-party evidence the series relies on — so the arguments rest on a common, citable base rather than re-arguing the facts each time — and is best read early, right after Paper 1. The Proposed FCC Rules examines the regulations the industry is treating as its answer, drawn entirely from the public record: what they require, who must carry them out, and when they take effect. It is best read after Paper 5, where the series turns from the coverage gap to the question of enforcement.

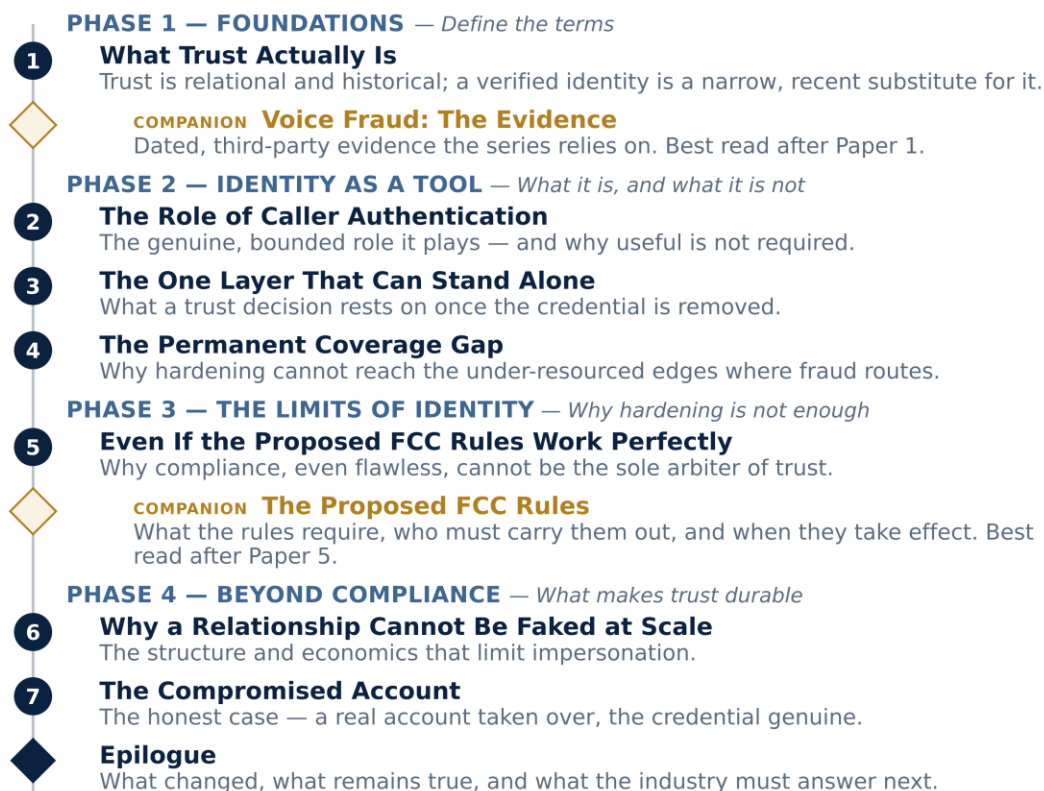
In reading order, then: Paper 1, followed by Voice Fraud: The Evidence, then Papers 2 through 5, then The Proposed FCC Rules, then Papers 6 and 7, and finally the Epilogue. The seven papers remain a single numbered argument; the companions sit where they are best read, not at the end, and the epilogue is read last.

The Trust in Communications series — in reading order

Seven papers, two companions that hold the evidence and the rules, and a concluding epilogue.

CENTRAL THESIS

Caller authentication is useful — but neither necessary nor sufficient for trust.



THE ARGUMENT ARC — HOW THE SERIES BUILDS THE CASE

1. Define

trust and identity as the starting terms

2. Understand

identity has a role, but it is bounded

3. Find the layer

the trust decision can rest here

4. Face the gap

authentication cannot reach every route

5. Examine compliance

even perfect rules are not the answer

6. Leverage economics

relationships are costly to fake at scale

7. Test the hardest case

compromise is real; divergence surfaces

HOW TO READ

In reading order: Paper 1, followed by Voice Fraud: The Evidence, then Papers 2-5, then The Proposed FCC Rules, then Papers 6-7, and finally the Epilogue. The seven papers form a single numbered argument. The companions sit where they are best read, not at the end. The epilogue is read last.

Read in order, or read the one that answers the question you came with.

Fig I-1 · series architecture and reading order

The Trust in Communications series — architecture and reading order: seven papers, with the two companions placed where they do the most work, and the concluding epilogue.

How the series is released

The series is released in sequence rather than all at once. The Introduction, Paper 1, and Voice Fraud: The Evidence anchor the opening — together they set the frame, make the core argument, and put

the shared evidence in the reader's hands, which the later papers draw on. The remaining papers follow in order, with The Proposed FCC Rules arriving alongside Paper 5, where the argument turns to enforcement. Each piece stands on its own as it appears; read together and in order, they form the single argument this introduction describes.

How this series is written

A word on posture. This series is written for readers who know this domain well, including those inclined to disagree. It holds to a discipline throughout: concede what is true, claim only what can be defended, and mark the genuinely unsettled as unsettled. It is an inquiry, not a sales sheet. It is published by ICA AI, Inc., which is building relationship-intelligence trust infrastructure for communications; its individual claims are meant to stand on the evidence and reasoning presented, not on the publisher's interest in the outcome. The fair account is the persuasive one.

Drafted with AI assistance; all analysis, claims, and conclusions rest with the series editor.

What Trust Actually Is

Why a Verified Identity Is Not the Same as Trust

TRUST IN COMMUNICATIONS · PAPER 1 OF 7 · ICA AI, Inc. · +Trusted Infrastructure · July 2026

A +Trusted Infrastructure series from ICA AI

Series editor: John Perkins, Chief Executive Officer

Modern communications security has quietly swapped one question for another. In place of “should I trust this?” it asks “is this identity verified?” — and treats a yes to the second as a yes to the first. This paper separates them. Trust is relational: earned across time, context, and repeated interaction. Identity verification is a point-check — valuable, but narrow, and blind to almost everything trust is made of. The distinction decides whether the systems being built to protect voice communications are aimed at the right target, and it is the premise on which every paper that follows depends.

The question we stopped asking

For most of human history, trust in communication was not a technical problem. You knew a voice because you had heard it a hundred times. You extended credit to a merchant because your family had traded with his for a generation. You believed the messenger because of who vouched for him and what he had done before. Trust was something that accumulated — slowly, in context, between specific parties who carried a shared history. It was never a stamp anyone wore; it was a state that grew between people who kept showing up. Relationship has long been the primary way humans have judged trust — refined across the generations that depended on it. The identity credential, by comparison, is a recent arrival.

What made it necessary was the telephone. A voice on a line carries none of the context a face or a shared history once supplied. The industry’s answer was to rebuild trust from a single substitute: identity. If we can prove who is calling, the reasoning goes, we can decide whether to trust the call. Over the last decade that reasoning has hardened into infrastructure — caller authentication, attestation, know-your-customer rules — all of it organized around verifying the identity behind a communication. It is a reasonable instinct. Identity is checkable, and the simplest scams work by faking or hiding who is calling. But the instinct quietly answered a narrower question in place of the one the industry had stopped asking out loud.

Because somewhere along the way a substitution happened: the question “should I trust this communication?” was replaced by the question “is this identity verified?” — and a yes to the second is now widely treated as a yes to the first. It is not. Identity and trust are different things, and conflating them is the quiet error underneath a great deal of what the industry is currently building. This paper is about that error, because everything downstream depends on getting it right.

What trust actually is

Trust is a relationship, not an attribute. It is not a stamp a party carries; it is a state that exists between two parties, and it is built from evidence that accumulates over time. Consider how you actually trust the people you trust. You trust your sister’s voice on the phone not because a credential — an authenticated caller ID — verified it, but because you have heard it your entire life — its cadence, its

turns of phrase, the things it would and would not say. Strip the history away and the credential alone would not be enough; keep the history and you can often recognize a caller before they have finished a sentence.

Three properties fall out of this that identity verification does not have:

- **Trust is historical.** It is a function of accumulated interaction. A relationship with a thousand prior exchanges is a different object than a first contact, even when the identity is equally certain in both. The relationship history is not an outgrowth of the trust; the trust is an outgrowth of the relationship history.
- **Trust is contextual.** The same verified party is trusted for some things and not others, at some times and not others, through some channels and not others. Trust is scoped to a relationship's established pattern, not granted wholesale, and a request that falls outside that pattern draws suspicion even from a party you know well.
- **Trust is reciprocal.** It lives on both sides. Each party holds its own record of the other, and those records have to agree. That two-sidedness is very hard to fabricate, precisely because it does not live in one place where a single forgery can reach it.

None of these properties can be read off a credential. They are properties of a relationship over time, and a point-in-time check has no access to any of them.

The call from your bank

A concrete example shows how much work the relationship is quietly doing. Your bank's fraud line calls. You trust the call partly because of the number it comes from — but far more because of everything the number, by itself, cannot tell you. The caller references a transaction you actually made. The exchange unfolds the way these exchanges always have. It comes at a time of day your bank actually calls. You know how this bank handles these calls as surely as it knows you — the details it raises, the rhythm the two of you have settled into over years of dealing with each other. Nothing in the request falls outside the pattern of every prior interaction you have had with them. The credential is one input among many, and not the decisive one; strip it away and the rest — the familiar pattern, the shared history, the fit with how your bank actually behaves — would still mostly carry the call.

Now change one thing: the caller asks you to move money to a new account, urgently, in a way your bank never has. The number still checks out. The identity is still verified. And your trust evaporates — not because the credential failed, but because the behavior left the relationship's established pattern. That single swing, from trust to suspicion with the credential held constant, is the whole argument of this paper in miniature. What moved was never the identity. It was the fit between the behavior and the relationship, which is the thing you were actually trusting all along.

What identity verification actually does

Identity verification answers exactly one question, and answers it well: who is the party presenting this communication? That is a genuinely valuable thing to know. Knowing origin deters casual abuse, raises the cost of the mass spoofing operations behind most robocalls, and makes calls traceable enough to hold the provider responsible. It would be a mistake to dismiss it, and this series does not. But notice the shape of what it provides. Identity verification is a point-check: it evaluates a claim at

a moment, returns a yes or no, and says nothing about the relationship, the context, or the history in which the communication sits. It is issued to a party, not built between parties. It is one-directional — a credential the caller carries, not a record the callee also holds. And once granted, it does not change with behavior: the same verification rides the first call and the thousandth, the honest call and the abusive one. These are not defects. They are the boundaries of what a point-check is. Identity verification tells you *who*. It does not, and cannot, tell you *whether* to trust.

Where the two come apart

The clearest place to see the gap is the case the industry finds most uncomfortable: a valid credential on a compromised account. When a real account is taken over, every identity check still passes. The credential is real. The verification is honest. And the communication should not be trusted at all — because the relationship it claims to belong to is not the one actually driving it. That single case is enough to establish the point. Identity verified, trust misplaced. If verifying identity were the same as establishing trust, that situation could not exist. It exists routinely, and it is the case a purely identity-based system is structurally unable to see.

The reverse gap is just as real, and just as instructive. You can trust a communication whose formal identity you never verified, because everything else — the timing, the knowledge, the pattern, the reciprocal record on your side — lines up with a relationship you already have. Humans have always done this, and are usually right to. The identity check is one input to a trust decision, sometimes a helpful one; it is not the decision itself. Both gaps point the same way: identity and the trust decision are different objects, and they come apart in exactly the cases that matter most for fraud. None of this implies that relationships are infallible — people routinely misplace trust. The question is whether a relationship carries information a credential alone cannot.

Identity vs. Trust — two different things

Identity answers “Who is calling?” Trust answers “Should I trust this?”

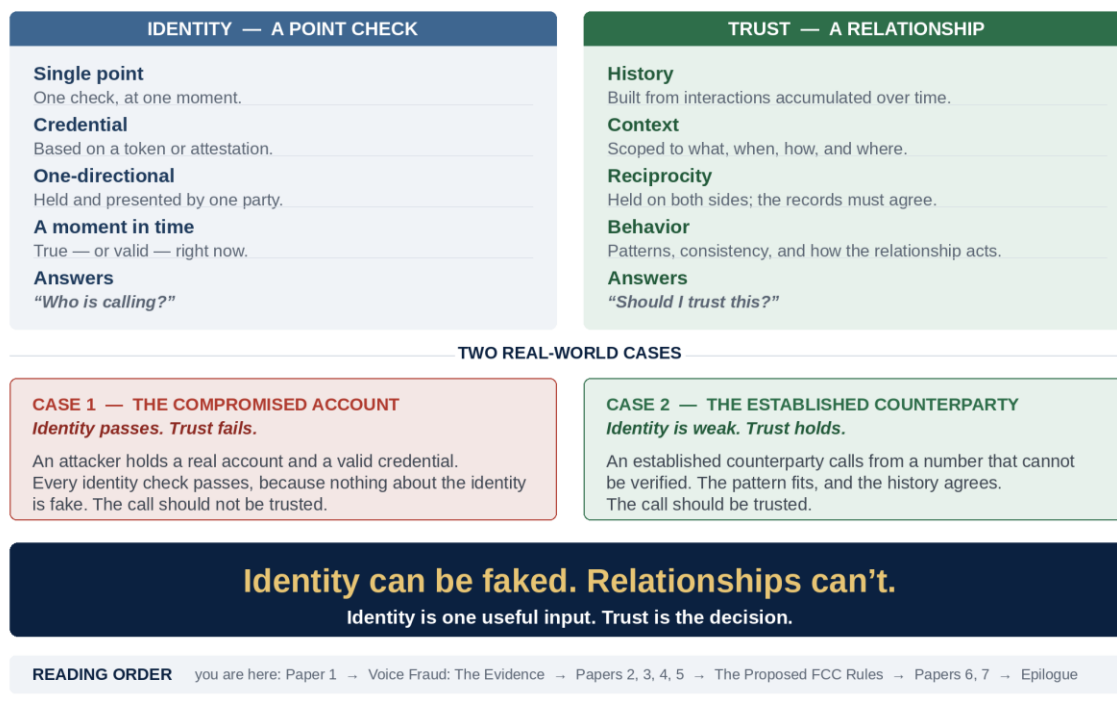


Fig 1-1 · conceptual — no data claimed

Identity vs. trust — two different questions, answered by two different kinds of evidence; the two cases below are where they pull apart.

Why this matters for voice

If identity and trust were interchangeable, the current direction of travel would be sound: verify harder, verify more, and trust follows. But because they are different things, an uncomfortable consequence follows. A system that verifies identity — however rigorously — has not yet made a trust decision. It has answered *who*. The question of *whether* still sits untouched, waiting for exactly the historical, contextual, reciprocal evidence a point-check was never built to carry. This matters more now, not less. Verifying a credential keeps getting cheaper and easier — and so does faking one. AI can now produce a convincing voice or face on demand. So the industry is leaning harder on identity at exactly the moment identity is becoming easier to counterfeit.

This is not an argument that identity verification should be abandoned. It is an argument that it should be seen for what it is: a useful input to a larger decision, not a substitute for it. The rest of this series follows that thread — what a trust decision actually requires, where identity helps and where it cannot reach, and what it would mean to build on the thing trust is actually made of, the relationship, rather than on the narrower thing we have been treating as its stand-in.

What this changes

Naming the distinction changes where the effort should go. If trust is relational and identity is a point-check, then pouring more attestation, more regulation, more verification into a stronger point-check improves the part of the problem that was already best handled and leaves the harder part — the relationship, the context, the behavior over time — exactly where it was. The productive move is not

to abandon the credential but to stop asking it to be the whole decision, and to build the layer that can actually weigh whether a communication fits the relationship it claims. That is the argument the series develops. This paper asks only that one premise be granted: identity and trust are not the same thing.

What this series claims, and what it does not

A note on posture, because it governs everything that follows. This series is written to be read by people who know this domain well, including those inclined to disagree. So it holds to a discipline: it concedes what is true, claims only what it can defend, and marks the genuinely unsettled as unsettled. It does not argue that caller authentication has failed — it has measurably reduced certain kinds of abuse. It does not argue that identity does not matter.

It argues something narrower and harder to refute: that a verified identity is useful, but not necessary and not sufficient for trust, because trust exists between two parties, while identity belongs to only one. Everything downstream — the evidence, the economics, the architecture, and the honest limits — rests on this one distinction. Get it right and the rest of the series follows. Blur it, and the entire debate about trust in communications is being conducted about the wrong thing.

Drafted with AI assistance; all analysis, claims, and conclusions rest with the series editor.

Voice Fraud: The Evidence

The third-party evidence for the negative case — that attestation is narrow, the threat has migrated, and the share of signed calls is uneven and, after five years, still a minority.

COMPANION VOLUME · ICA AI, Inc. · +Trusted Infrastructure · July 2026

A +Trusted Infrastructure series from ICA AI

Series editor: John Perkins, Chief Executive Officer

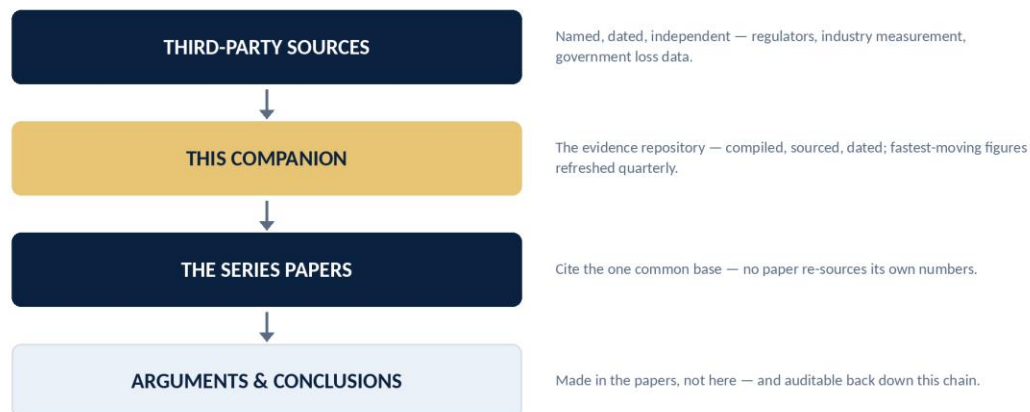
Every figure in this volume carries a named third-party source and a date; the fastest-moving figures are refreshed each quarter.

How to use this document

This is the reservoir. It carries no argument of its own; it holds the dated, sourced facts that every other paper in the series draws from, so the arguing papers can cite one common evidence base instead of each re-sourcing the same numbers. Read it as a reference, not a narrative. It is the canonical evidence repository for the series: when a statistic changes, it changes here first, and the papers inherit the update through their citations.

How This Evidence Is Used

The audit path from source to conclusion — evidence is compiled here; arguments are made in the papers.



To audit a conclusion: walk the chain in reverse — conclusion → argument → citation → source.

Fig E-1 · how this evidence is used

Read this first: it is the standing rule for the companion. Every figure that follows resolves to a named, dated, third-party source, and every paper cites this shared base rather than re-sourcing its own numbers.

A note on scope — This document gathers the *negative case* only: the third-party, dated evidence that verifying identity is not sufficient for trust — that attestation is narrow, that the threat has migrated past it, and that the share of signed calls is uneven and still short of half after five years of mandate. That case is built entirely from independent sources and is fully citable. The *positive case* — the claim that relationship-based trust outperforms identity — is deliberately excluded, because it is ICA AI’s own

thesis, argued in the papers and not yet independently established. Keeping the two apart is what lets everything here stand as fact anyone can check.

A discipline note that runs through every section: where a trend has several causes, this document says so. The measured drop in spoofing over recent years was driven by enforcement, call analytics, and labeling rolled out at the same time as attestation — attestation earned part of that win, not all of it. Crediting it honestly is what keeps the rest of this document credible.

1. Attestation coverage — what is actually signed

STIR/SHAKEN signs a minority of the calls that reach American phones. As of June 2026, signed calls were 48.8% of all calls observed at termination — a 20-month high, recovering toward the October 2024 peak as the all-IP transition retires the segments that strip signatures — with the highest attestation level (A) accounting for 30.8% and B for 4.4%. Five years after the mandate, fewer than one call in three arrives with the full-confidence attestation the framework was built to produce, and the overall share has yet to cross half. The trajectory from the same source’s monthly statistics: roughly 24% at year-end 2021 — stuck there since August 2021 — 26.2% at year-end 2022, 33% at year-end 2023, a 49.3% peak in October 2024 falling to 47.0% by December, a slide to 42.3% by April 2026, and a recovery to 48.8% in June (TransNexus monthly STIR/SHAKEN statistics, 2021–2026). The recovery is transport, not trust: what is rising is the share of calls whose signature survives the path — the grade distribution, and who holds the top grade, are unchanged.

Metric (terminating traffic)	Value	As of	Source
All signed calls	48.8%	Jun 2026	TransNexus
A-level (full attestation)	30.8%	Jun 2026	TransNexus
B-level (partial)	4.4%	Jun 2026	TransNexus
C-level (gateway)	7.8%	Jun 2026	TransNexus
Signed calls arriving with signature intact	44.5%	Feb 2026	TransNexus / ACA Int'l
Top-10 prolific robocall signers' calls at A-level	97.4%	Jun 2026	TransNexus / ACA Int'l

Two figures in that table matter more than the headline coverage number, and both cut against attestation as a trust solution:

The signature often gets stripped in transit. In February 2026, only 44.5% of calls that were signed at origin still carried an intact signature when they reached the destination. Roughly half the verification is lost in transit, principally where calls cross older non-IP (TDM) network segments that strip the SIP Identity header (documented as a mechanism in Section 5). A trust signal that disappears mid-path for half of traffic cannot be the layer trust rests on.

How 100% Becomes 48.8%

From all calls to signed-at-termination — two loss stages, one signed minority remains.

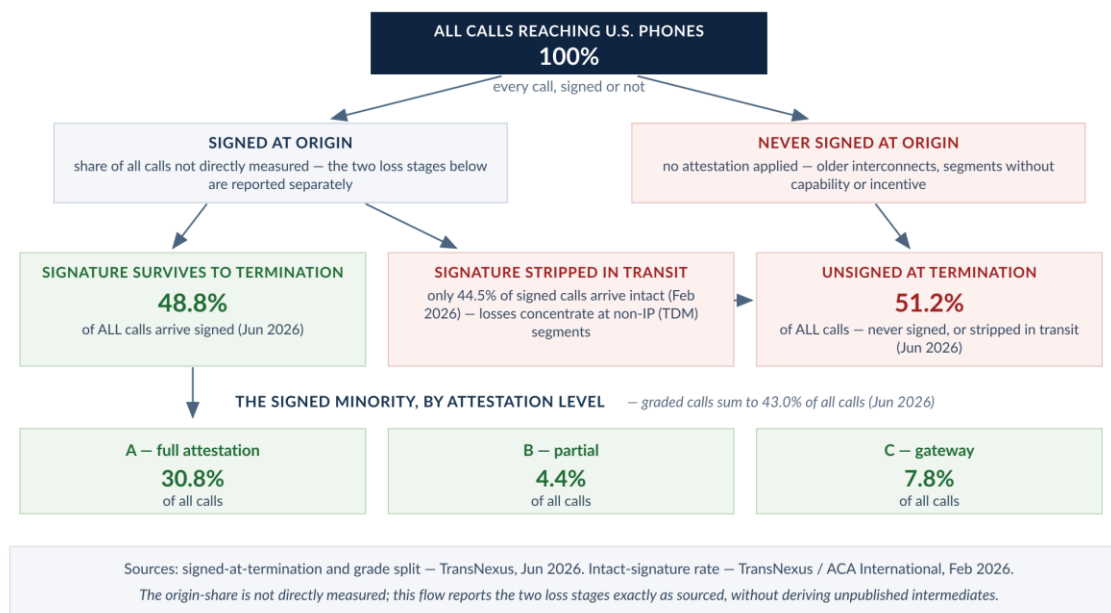


Fig E-2 · how 100% becomes 48.8%

The same 100%, twice divided — once at the origin, once in transit.

The legacy segments where this happens are not a rounding error. Among smaller carriers, only 17.5% of the call traffic between them was signed in 2025 (TNS, 2026 Robocall Investigation Report). About 83% of rural carriers (NTCA members) report at least some IP switching in their networks — leaving a real minority with none (NTCA comments to the FCC, 2025). And even on Tier-1 networks, roughly 40% of traffic in places lacks a signature because of TDM bypasses (TNS). Each figure measures a different slice, but together they show how much of the network still cannot carry an authenticated signature.

Fraudsters hold A-level attestation. Among the ten most prolific robocall-signing providers, 97.4% of calls carried A-level attestation — the highest trust level the framework issues. Attestation certifies that a provider vouches for the caller's right to use a number; it does not certify that the call is wanted, honest, or from who the recipient thinks. High-volume robocallers routinely obtain the top grade. This is the single most important fact about attestation as a trust signal: it is present and green on calls the recipient would never want.

Framing discipline — This is narrow-and-insufficient, not failure. Attestation measurably reduced number spoofing and is a genuine, useful hardening of the origin layer. The point is not that it broke; it is that it grades where a call entered the network — not who is behind it, and not the trustworthiness of the relationship behind it — and it is issued at the top level to the exact actors the recipient wants screened out.

2. Volume — the unwanted-call picture

Total robocall volume has plateaued while the composition has shifted toward the calls people least want. U.S. consumers received 52.5 billion robocalls in 2025 — essentially flat against 52.8 billion in 2024 (down less than 1%). But within that flat total, unwanted telemarketing and scam calls rose to 29.6

billion, or 57% of all robocalls, up from 49% the year before — a 15.4% year-over-year increase in the worst category even as the overall number held steady. Notifications and payment reminders (the benign categories) fell.

Robocall composition	2025 volume	Share	YoY	Source
Total robocalls	52.5B	100%	-0.6%	YouMail
Telemarketing + scam	29.6B	57%	+15.4%	YouMail
Notifications	13.3B	25%	-13.5%	YouMail
Payment reminders	9.4B	18%	-19.6%	YouMail

The story the volume data tells is not that calling is being cleaned up. It is that the mix is getting worse: the wanted, informational calls are receding and the unwanted, higher-risk calls are taking their place — during the same period attestation coverage stopped growing.

3. Fraud losses — the government-sourced thesis

This is an important line of evidence in the series, because it comes from a neutral government source and points in the opposite direction from attestation coverage. Reported U.S. fraud losses have risen every year while attestation was rolled out and hardened, reaching \$12.5 billion in 2024 and a record of approximately \$16 billion in 2025 — a 25% jump in a single year.

Reported U.S. fraud losses (all channels)	Amount	Source
2021	\$5.8B	FTC Consumer Sentinel Data Book
2022	\$8.8B	FTC Consumer Sentinel Data Book
2023	\$10.0B	FTC Consumer Sentinel Data Book
2024	\$12.5B (+25%)	FTC Consumer Sentinel Data Book 2024 (rel. Mar 2025)
2025	≈\$16B (+~25%)	FTC, congressional testimony (Apr 2026)

The 2024 jump was not an artifact of more people complaining. The number of reports was roughly stable; what rose was the share of reporters who actually lost money — from 27% in 2023 to 38% in 2024. More of the fraud that was attempted succeeded.

Where the phone sits. Across contact methods, email is the #1 method fraudsters use to reach victims, phone calls are #2, and text messages #3. Investment scams led all categories at \$5.7 billion, with imposter scams second at \$2.95 billion (government-imposter alone \$789 million).

On the fraud-loss figure — The \$12.5 billion covers all fraud across every channel, not phone alone. From a neutral, government source, attestation coverage rose over the same years that reported fraud losses rose — consistent with the losses migrating to vectors attestation cannot see.

Five Years of Attestation, Still Under Half; Losses Only Climbed

Two measurements, two units, one shared timeline — separate panels so neither implies a quantitative relationship with the other.

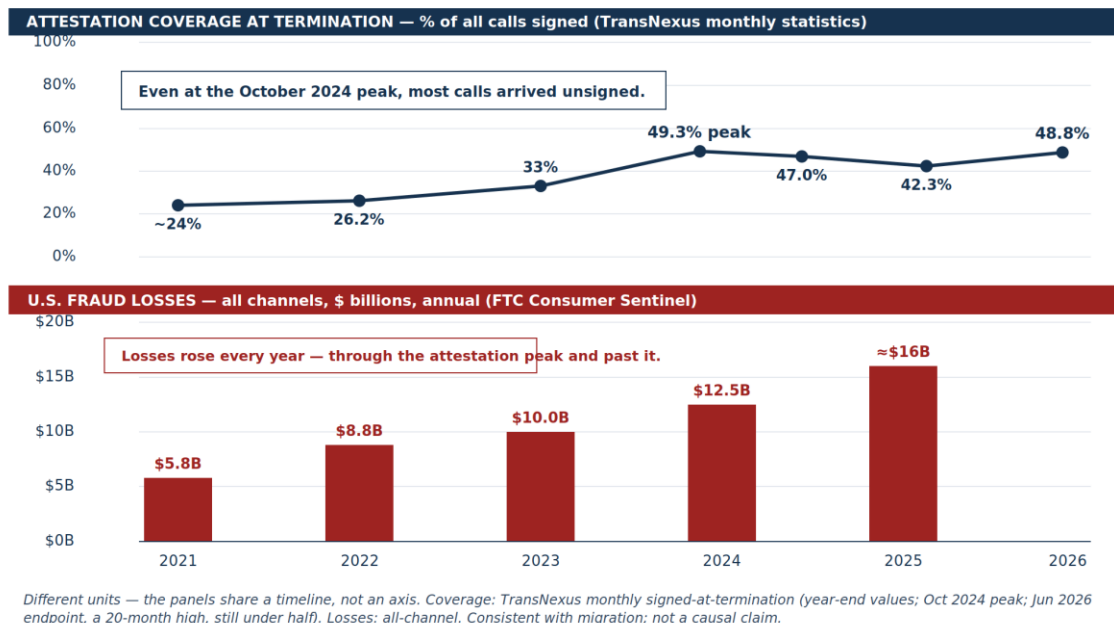


Fig E-3 · attestation coverage vs. fraud losses

The figures behind the two panels: coverage 49.3% at its October 2024 peak, 42.3% by April 2026, recovering to 48.8% in June 2026 — still under half after five years; reported losses a record ≈\$16B in 2025. Sources: TransNexus (signed-at-termination); FTC Consumer Sentinel (all-channel).

4. Threat migration — why the threat outran the standard

Attestation answers a question about the origin of a call. The threat has moved to a question attestation never asked: is the voice on the line who it claims to be? AI voice synthesis has turned that from a corner case into a mainstream attack in roughly two years, and the third-party measurements are stark.

The Threat Outran the Standard

Attestation answers where a call entered the network. In roughly two years — 2024 to 2026 — the threat moved to a question it never asked: is the voice real?

THE SUPPLY SIDE	attacker capability	THE DEMAND SIDE	consumer exposure
+1,300%	deepfake fraud attempts in 2024, measured against 2023 — from analysis of 1.2 billion calls	1 in 4	Americans say they received a deepfake voice call in the past 12 months
7 / day	deepfake calls flagged per customer by the end of 2024 — against roughly one per month across the entire customer base before ChatGPT	24%	another quarter are not sure they could tell a deepfake from a real call — together, nearly half the population
every 46 sec	a fraud attempt in U.S. contact centers, on average	~2 to 1	asked who is winning, Americans say scammers are beating the carriers
+173%	rise in synthetic-voice calls, Q1 → Q4 2024	9.9 / week	unwanted calls per person, +16% compounded since 2023
\$44.5B	projected contact-center fraud exposure for 2025 — a forecast, not a measurement		
Pindrop, Voice Intelligence & Security Report — June 12, 2025		Hiya, State of the Call 2026 — March 2, 2026. Survey of 12,000+ consumers across six markets; U.S. figures.	

WHY THIS IS MIGRATION, NOT NOISE

- Every marker above can, in principle, ride a properly attested call.
- Attestation is unrelated to voice authenticity — the two are independent, and neither tells you anything about the other: a correctly signed A-level call can carry a cloned voice.
- The standard hardened the origin; the threat moved to the content — the one surface no standard reads.

Note on the source. Pindrop's release pairs the 1,300% year-over-year rate with the "one per month to seven per day" frequency as though the two measured the same thing. They do not: the earlier figure counts deepfake calls across Pindrop's entire customer base per month; the later counts them per customer per day. They are reported separately above.

Fig E-4 · threat migration — dated third-party markers

Supply and demand, measured independently by two parties. Every marker here describes a threat a correctly signed call carries untouched: attacker capability (Pindrop, Jun 2025) and consumer exposure (Hiya, Mar 2026).

The supply side — synthetic voice at scale

Pindrop's 2025 Voice Intelligence & Security Report (released June 12, 2025) recorded a **1,300% increase in deepfake fraud** over 2024, measured across 1.2 billion customer calls, and a 173% rise in synthetic-voice calls between the first and last quarter of 2024. Pindrop separately reports flagging roughly seven deepfake calls per customer per day by the end of 2024, against about one per month across its entire customer base beforehand. Pindrop's own release quotes those two figures together as though they measured the same thing. They do not — the bases differ — and they are reported apart here. In U.S. contact centers, a fraud attempt now occurs on average every 46 seconds. Synthetic-voice attacks rose 475% in insurance, 149% in banking, and 107% in retail in a single year, and Pindrop put 2025 contact-center fraud exposure at \$44.5 billion.

The demand side — what consumers are already seeing

Hiya's State of the Call 2026 (released March 2, 2026) found that **1 in 4 Americans received an AI deepfake voice call in the prior twelve months**, and a further 24% are not sure they could tell a deepfake from a real call — so nearly half the population has either encountered AI voice fraud or cannot reliably distinguish it. The 24% is uncertainty, not measured failure: Hiya's wording is that they "aren't sure they could tell the difference." Asked who is winning the fight between carriers and scammers, consumers picked the scammers by nearly two to one. Unwanted calls now run 9.9 per week per person and have grown 16% compounded annually since 2023; 38% of subscribers say they would switch providers over inadequate AI-scam protection.

Threat-migration marker	Value	Source (dated)
Deepfake fraud growth, 2024	+1,300%	Pindrop, Jun 2025
Deepfake-fraud frequency (different base)	~1/month, whole customer base → ~7/day, per customer	Pindrop, Jun 2025
Contact-center fraud interval	every 46 sec	Pindrop, Jun 2025
Contact-center fraud exposure (2025, projected)	\$44.5B	Pindrop, Jun 2025
Americans who got a deepfake voice call (12 mo.)	1 in 4	Hiya, Mar 2026
Not sure they could distinguish deepfake from real	24%	Hiya, Mar 2026
Consumers say scammers are winning vs. carriers	~2 to 1	Hiya, Mar 2026

The common thread: every one of these attacks can, in principle, ride a properly attested call. Number attestation and voice authenticity are independent — neither tells you anything about the other. A correctly signed A-level call can carry a cloned voice. The standard hardened the origin while the threat moved to the content and the relationship.

5. The mid-path laundering mechanism (documented as fact)

When a call crosses from IP to legacy TDM and back to IP, the SIP Identity header that carries the STIR/SHAKEN signature is stripped at the TDM segment. The call then re-enters an IP network unsigned and is re-attested downstream — often at the lower C level, but per that same account sometimes re-signed at A by the re-originating provider. Deliberate header stripping by rogue or loosely regulated intermediate providers produces the same result on purpose.

The measured consequence is in Section 1: only 44.5% of signed calls arrive with their signature intact (February 2026). A signal that is silently regenerated mid-path — and can be upgraded to full attestation after the original chain of custody is broken — carries information about the last hop, not about the caller. A signal with those properties cannot be load-bearing for the trust decision.

6. Regulatory status — the backbone

The regulatory record here was verified against primary sources; this section summarizes it and re-confirms the fast-moving dockets as of this document's date. The single most important discipline: distinguish adopted rules from proposals, and origin-level obligations from relationship-level ones.

Regime	Status	What it actually is (and the docket)
KYC — point-in-time	Adopted	Providers carry an ongoing know-your-customer / know-your-upstream duty via the robocall-mitigation certification regime, enforced by RMD removal and monetary settlements. Real, and origin-level.
KYC — 'continuous'	Proposed	Continuous monitoring is a pending FNPRM (FCC 26-27, CG 17-59; companion know-your-upstream item FCC 26-32, WC 17-97). Considered at Spring 2026 Open Meetings; no codified continuous-monitoring rule, no effective date.
Org-ID / RCD	Real, industry-built	Delegate certificates (ATIS-1000092), Rich Call Data (RFC 9795), and CTIA Branded Calling ID ride the STIR/SHAKEN

Regime	Status	What it actually is (and the docket)
		certificate system. NOT an FCC mandate; the only live FCC action is a pending caller-ID display item (FCC 25-76).
Proof-of-Life / human-bot ID	No docket	No FCC or FTC proceeding proposes a human-vs-bot caller credential. CG 23-362 proposes disclosure of AI-generated calls only; the FTC Voice Cloning Challenge was a one-time prize competition. No ATIS or IETF work item found.

The enforcement record — how removal has actually been used. The expanded certification requirements set a compliance deadline of February 26, 2024. After more than 2,000 providers failed to meet it, the Enforcement Bureau issued a show-cause order in December 2024 directing 2,411 providers to cure their filings or face removal. On August 6, 2025, the Bureau removed an initial 185 providers from the database; on August 25, 2025, it removed the certifications of more than 1,200 more — the first mass sweep in the database's history. On January 2, 2026, it removed 1,203 more. Removed providers may not refile without prior approval from the Enforcement and Wireline Competition Bureaus. From missed deadline to disconnection: eighteen months. (FCC Enforcement Bureau, Aug 6 and Aug 25, 2025.)

The distinction that matters. Even at its most expansive proposed form, the regulation operates at the origin/provider level — a carrier vetting who is allowed to originate traffic. Nothing adopted or proposed reaches the pair-level relationship layer — who actually trusts whom. The regulatory record leaves the category question entirely intact, which is itself a finding: neither pole can claim the FCC as an ally on the category.

7. What this evidence supports — and what it does not

Stated plainly, so no paper built on this reservoir overreaches:

Supported, and fully citable (the negative case):

Attestation signs a minority of calls (48.8%) and has yet to cross half after five years of mandate; roughly half of signatures are lost in transit; the highest attestation level is routinely held by prolific robocallers; unwanted calls are rising as a share of a flat total; reported fraud losses hit a record of approximately \$16B in 2025 (up ~25% over \$12.5B in 2024) with a rising success rate; and the threat has migrated to AI voice, which rides attested calls untouched. All third-party, all dated.

NOT supported here, and deliberately excluded (the positive case):

That relationship-based tracking outperforms identity is a claim this document does not make, because it is not yet external, third-party research — it would be ICA's own claim to prove, not an independent finding. Until it is, it stays out of this evidence base. Keeping it out is what makes everything above trustworthy.

Bottom line — The evidence cleanly establishes that attestation is narrow, leaky in transit, blind to voice authenticity, and issued to bad actors — and that the threat has moved past it while losses climbed. That is the whole burden this document needs to carry. The case that relationship-based trust does better is made elsewhere, in the papers, and rests on its own separate evidence.

The Boundary of Evidence

What this companion establishes — and where it deliberately stops.

ESTABLISHED BY THIS DOCUMENT — third-party, dated

COVERAGE	A minority of calls is signed — 48.8% at termination (Jun 2026), a 20-month high after a 49.3% peak (Oct 2024) and a slide back — and still short of half after five years.
TRANSIT	Roughly half of signatures are lost in transit — only 44.5% of signed calls arrive intact (Feb 2026).
ISSUANCE	The top attestation grade rides prolific robocall traffic — 97.4% of the top-10 signers' calls carry A-level (Jun 2026).
VOLUME	Unwanted calls are rising inside a flat total — scam and telemarketing reached 57% of robocalls (+15.4%) while overall volume stayed level (2025).
LOSSES & THREAT	Reported U.S. fraud losses rose every year of the rollout — \$5.8B (2021) to ~\$16B (2025) — while the threat migrated to voice authenticity.
REGULATION	The rules are real but proposed and phased — and the remedy is an enforcement sequence, not a switch.

THE BOUNDARY OF EVIDENCE — THIS DOCUMENT STOPS HERE

NOT ESTABLISHED BY THIS DOCUMENT

That relationship-based trust does better. That is the series' own claim — argued in the papers, on its own evidence, and deliberately excluded here.

Fig E-5 · the boundary of evidence

The boundary of evidence — what this companion establishes, and the claim it deliberately does not make.

Sources & verification log

All figures below are current as of the dates shown. Coverage and threat figures move within a quarter; the dated dockets and monthly statistics are the authoritative current source.

#	Figure(s)	Source	Dated
1	Signed 48.8%; A 30.8% / B 4.4%	TransNexus, STIR/SHAKEN statistics	Jun 2026
2	44.5% signatures intact (Feb 2026); 97.4% top signers at A (Jun 2026)	TransNexus via ACA International	Feb / Jun 2026
3	52.5B robocalls; 29.6B (57%) unwanted	YouMail Robocall Index	2025 (Jan 2026)
4	Fraud losses \$5.8B→\$8.8B→\$10B→\$12.5B→\$16B (2025); 27%→38%	FTC Consumer Sentinel Data Book 2024; FTC testimony (2025 figure)	Mar 2025 / Apr 2026
5	Email #1 / phone #2 / text #3; investment \$5.7B; imposter \$2.95B	FTC Consumer Sentinel Data Book 2024	Mar 2025
6	Deepfake fraud +1,300% (1.2B calls); ~7/day per customer; \$44.5B projected	Pindrop Voice Intelligence & Security Report	Jun 2025
7	1-in-4 deepfake call; ~2:1 scammers winning; 9.9/wk	Hiya, State of the Call 2026	Mar 2026
8	KYC/KYUP FNPRMs; Org-ID rails; Proof-of-Life none	FCC 26-27 / 26-32 / 25-76; CG 23-362	Jul 2026

Primary sources: transnexus.com/blog/2026/shaken-statistics-April; acainternational.org (STIR/SHAKEN Feb 2026); YouMail Robocall Index 2025 year-end release; ftc.gov Consumer Sentinel Network Data Book 2024 (and the March 2025 press release) — the 2025 calendar-year figure comes from FTC congressional testimony of April 2026, not from a 2025 Data Book, which had not been published as of this writing; Pindrop 2025 Voice Intelligence & Security Report; hiya.com

State of the Call 2026; docs.fcc.gov FCC-26-27A1 and related dockets; FCC Enforcement Bureau RMD removal orders — DOC-413519A1 (Aug 6, 2025) and DOC-414073A1 (Aug 25, 2025); further removal order, Jan 2, 2026.

The Role of Caller Authentication in Establishing Trust

Useful, but Not Necessary

TRUST IN COMMUNICATIONS · PAPER 2 OF 7 · ICA AI, Inc. · +Trusted Infrastructure · July 2026

A +Trusted Infrastructure series from ICA AI

Series editor: John Perkins, Chief Executive Officer

Caller authentication earns a real place in a trust decision. At the moment two parties share no history, an attestation signal is often the most useful thing available. But useful is not the same as required. Attestation is a fixed credential that does not change as a relationship does; the relationship signal is designed to update with every interaction. The credential is frozen; the relationship learns. A single relationship, followed from its first call to its hundredth, draws the line the series turns on: authentication is useful, and most useful at the beginning, but it is not necessary, because the thing that ultimately carries trust is the relationship, not the identity credential.

What authentication genuinely provides

At genuine cold start — the first contact between two parties with no shared record — a caller-authentication signal is often the best thing a receiving system has to go on. In the United States that signal is STIR/SHAKEN, the framework that attaches a cryptographic attestation to a call as it enters the network; where this series says the standard, that is the standard it means. It arrives before the call connects: it is a pre-ring signal, present at the exact moment when nothing behavioral yet exists. Remove it at that moment and you are left with less at the decision point, not more. So the honest starting position is not that authentication is weak or that the last decade of work on it was misdirected. It is that authentication is valuable, and most valuable precisely when a relationship is youngest and there is nothing else to weigh.

A pharmacy the customer has never used calls for the first time about a transferred prescription. There is no history to draw on, no pattern to match, no reciprocal record on either side. In that moment a validly attested number is a genuine signal — a caller cannot produce one just by spoofing a number for free — and it gives the receiving side something real to weigh where it would otherwise have nothing. The argument that follows is not that the signal is useless. It is that the signal is temporary, and that its value is highest at exactly the point the relationship is weakest.

What the same evidence base also shows is where that value stops. The top attestation grade is routinely carried by prolific robocallers — among the ten most prolific robocall-signing providers, 97.4 percent of calls arrive at A-level, the highest grade the framework issues — because the grade certifies how a call got onto the network — which provider originated it and vouched for the number — not whether the call is worth trusting. And attestation and voice authenticity are independent — neither tells you anything about the other: a correctly signed, top-grade call can carry a cloned voice, because the credential vouches for a number's provenance, not for who is actually speaking. The signal that is genuinely useful at cold start is, demonstrably, not a trust decision on its own.

Frozen by design

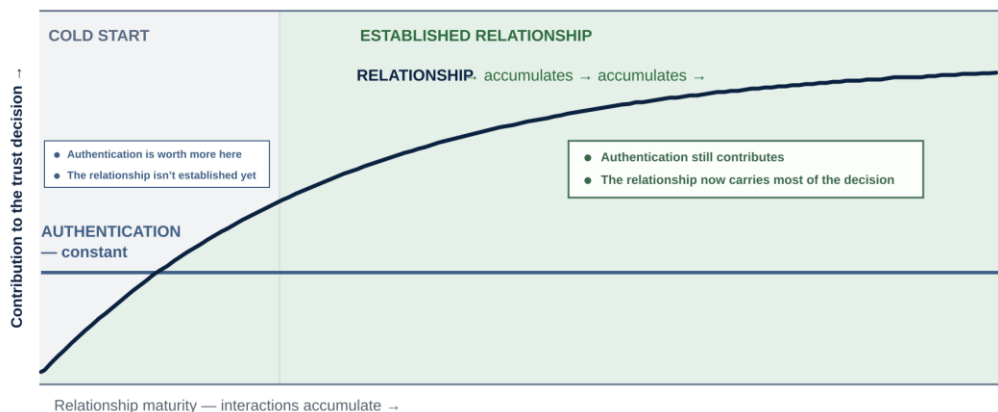
The defining property of a credential is that it does not change with behavior. The attestation on call one is the attestation on call one thousand; the honest call and the abusive one carry the same grade. This is not a defect — it is what a credential is, and what makes it verifiable at all. But it has a consequence that is easy to miss. Authentication says exactly the same thing about a specific pair of parties on their first interaction and their hundredth. It cannot grow more accurate about that pair over time, because it was never built to observe the pair. It authenticates a party in isolation, once, and issues the same verdict forever — which is precisely why it is strongest when there is no history to be more accurate about, and why its relative value falls as history accumulates around it.

The signal that learns

A relationship signal is defined by the opposite property: it changes. The relationship layer is designed to update its picture of a specific pair with every genuine interaction — timing, cadence, the knowledge exchanged, the reciprocal record each side holds — so that the signal available at the hundredth call is richer than the one available at the first.

Relative Contribution to a Trust Decision

The credential never gets worse. The relationship gets better — and the decision migrates toward the signal that accumulates.



One signal is frozen at issuance. The other is designed to update with every interaction.

Conceptual — the curves are illustrative shapes, not measured data.

Fig 2-1 · relative contribution to a trust decision

The credential is worth the same on the hundredth call as on the first; the relationship signal is designed to update with every interaction.
Conceptual — the axes carry no scale.

A fair objection: number-reputation systems already learn over time, which makes the word frozen sound too strong. The answer is what each system learns about. Reputation systems learn about a number — a single party in isolation — building a score that is still issued to that party and carried by it wherever it goes. The relationship signal learns about a pair of parties: not how a number behaves in general, but whether this particular communication fits the history these two specific parties share. A number can carry a spotless reputation into a relationship it has no business being in, and a number-level system will wave it through; the pair-level signal is the thing built to notice.

That is the asymmetry that matters — not whether a system learns, but whether it learns about a party or about a relationship.

There is a second distinction the industry tends to blur, and it is load-bearing too. Attestation is a pre-ring signal, present before the call connects. The recipient's own decision to block, screen, or trust is post-ring feedback, available only afterward. They are not substitutes. At true cold start there is no post-ring history yet — which is exactly why authentication is most valuable there. As interactions accumulate, the behavioral record grows and the frozen pre-ring credential becomes a smaller share of what the decision actually rests on. Return to the pharmacy. By the twentieth genuine call, the receiving side has a texture to match against — the hours it calls, the systems it references, the way its requests are shaped — and the attestation that carried the first call is now one small input among many. The credential did not get worse. The relationship simply got better, and the decision migrated toward it.

Same Credential, Different Relationship

The pharmacy, twice — a yes/no check on every call, and the history that does or does not stand behind it.

The diagram illustrates two authentication scenarios side-by-side, separated by a vertical dashed line.

CALL 1 — first contact (Left side, light blue background):

- Text:** "The check stamps. Nothing stands behind it yet."
- Authentication:** A box labeled "AUTHENTICATION — a yes/no check" contains a single green checkmark icon and the text "yes — verified, this call".
- Relationship:** A box labeled "RELATIONSHIP" contains the text "none yet — no history to read".
- Summary:** "The signed credential is the best signal going — the only signal on the board."

CALL 20 — established (Right side, light green background):

- Text:** "The same check stamps — and this time it isn't alone."
- Authentication:** A box labeled "AUTHENTICATION — a yes/no check" contains a single green checkmark icon and the text "yes — verified, this call".
- Relationship:** A box labeled "RELATIONSHIP — the calls remembered" contains a horizontal row of 20 green checkmark icons.
- Summary:** "Hours, systems, request shapes — the credential is now one signal among twenty."

Bottom Summary: A dark blue bar at the bottom contains the text: "Authentication is the same yes/no on every call. A relationship is every prior call, remembered."

Conceptual — glyphs are illustrative; the count marks the nineteen prior calls in the example.

Fig 2-2 : same credential, different relationship

Same credential, different relationship — the yes/no check is identical on both calls: nineteen remembered calls now stand behind the second.

What “supersede” means — and what it does not

As a relationship matures, the behavioral signal is designed to outweigh the credential in the trust score. That is a precise claim, and worth stating precisely, because it is easy to overstate into something this series does not mean. Supersede does not mean discard. Consider a top-grade call — validly signed, highest attestation — that the recipient blocks because everything about it is wrong. That event should lower trust for that pair and teach the graph, even though the credential was real and correctly issued. The credential remains an input; it simply stops being the deciding input once there is an established relationship to weigh it against. How the engine weights these inputs against one another over time is a matter of implementation; the structural point does not depend on it. So is where the layer runs: whether the relationship signal is computed by a carrier, an endpoint, or another trusted system is outside the scope of this paper. One signal is frozen and one learns, and over a relationship's life the one that learns comes to carry more of the weight.

Scaffolding, not floor

The accurate picture, then, is that authentication is scaffolding. It is most load-bearing when the structure is new and carries less as the structure sets — seemingly necessary at the start, superseded as the relationship signal matures. Scaffolding is not a criticism; you cannot raise the building without it, and pretending otherwise would be its own kind of overclaim. But no one confuses the scaffolding with the wall. This is also the honest answer to the strongest objection: that you cannot build trust on an unauthenticated network. At cold start that is true — where there is no history, authentication is the best signal going. What is not true is that this holds forever. The claim is not that authentication is dispensable in general; it is that it is temporary in any given relationship — the support a relationship needs most at the start, and less and less as it grows able to stand on its own.

What this changes in practice

If authentication is scaffolding rather than floor, the practical consequence is about where to invest. The instinct of the last decade — and of the proposed rules being written now — has been to make the credential stronger: harder to forge, more widely mandated, more finely graded. That work has value, and this paper does not argue against it. But it runs into a ceiling the later papers document, because a stronger credential still says the same thing about the thousandth call as the first. The layer with more headroom is the one that improves with every interaction: the relationship. Investing there does not mean abandoning authentication; it means recognizing which layer is scaffolding and which is structure, and putting the weight on the part that is designed to keep getting better after the credential has said all it can say.

There is a second consequence, easy to overlook. Because authentication is most valuable at cold start and least valuable once a relationship matures, its value is concentrated in a shrinking share of interactions over time. Every relationship that establishes itself is one more relationship for which the credential has become a minor input — while the credential's cost, in mandates and infrastructure and grading rigor, is paid on every call regardless. The economics of the two layers run in opposite directions: authentication continues to require the same infrastructure investment even as it contributes less with accumulating history, while the relationship signal costs its effort up front and contributes more as history accumulates. A trust system that understands this does not choose between them; it leans on authentication exactly where it is strong, at the beginning, and lets the relationship carry the growing remainder — which is most of the interactions, most of the time.

Useful, but not necessary

The two words sit side by side, and the whole series balances on the space between them. Useful: yes, plainly, and most of all at the beginning of a relationship. Necessary: no — because the thing that comes to carry a trust decision is the relationship itself, and a relationship can be built, observed, and weighed with authentication as an accelerant rather than a prerequisite. That distinction sets up the harder question the next paper takes on directly. If caller authentication were removed from the architecture entirely, would a trust decision still be made — and if so, which layer was doing the real work all along?

Drafted with AI assistance; all analysis, claims, and conclusions rest with the series editor.

The One Layer That Can Stand Alone

What Communication Trust Actually Rests On

TRUST IN COMMUNICATIONS · PAPER 3 OF 7 · ICA AI, Inc. · +Trusted Infrastructure · July 2026

A +Trusted Infrastructure series from ICA AI

Series editor: John Perkins, Chief Executive Officer

Remove caller authentication from the picture entirely and ask a blunt question: can a trust decision still be made another way? This paper argues it can — that the relationship layer is the one layer able to stand alone, while the identity layer, however well verified, cannot render a trust decision by itself. That makes authentication an input that looks load-bearing but is ultimately optional, and makes the relationship the thing trust actually rests on.

A cautionary case: false certainty is worse than none

In 2002 Arthur Andersen, one of the most trusted names in auditing, collapsed within months. What destroyed the firm was not that its certifications were missing — they were present, signed, and authoritative. They were simply wrong — and they were acted upon precisely because they looked authoritative. A missing signal makes people check. A confident, valid-looking signal makes people act — no one re-verifies what looks certified. That is how Andersen's wrong certifications did their damage. Auditing is not telecommunications, and the analogy claims no more than one shared structural property: a signal acted on because it looks authoritative, regardless of what stands behind it. The lesson for any trust system is uncomfortable: a certification that is trusted because it looks authoritative is more dangerous, not less, when it is wrong. Any architecture that treats a valid identity credential as the trust decision inherits exactly that failure mode — it converts a checkable claim into an unquestioned one. A stolen credential passes that check perfectly, and the system reports its highest confidence at exactly the moment it is most wrong.

The load-bearing test

There is a simple structural test for whether a component is load-bearing or optional, borrowed from the way engineers talk about buildings: remove it and see if the structure still stands. Apply it to caller authentication. Remove authentication and a relationship-based trust decision can still be rendered — from history, context, timing, the knowledge two parties share, and the reciprocal record each holds of the other. The decision is weaker at cold start — the very first contact, before any history exists — and this paper will not pretend otherwise, but it is still a decision, and it is still made on evidence.

Now run the test the other way, because a test you only run in one direction proves nothing. Remove the relationship and keep only the credential, and what remains is a decision of the bluntest kind — binary: pass anything validly signed, block anything unsigned. That is a real decision — identity does decide something. But look at what it decides. Roughly 49% of calls arrive signed and 51% arrive unsigned (June 2026). Applied literally, the credential-only gate blocks half of all calls, the

overwhelming share of them legitimate. No carrier does that, which is the tell: the unsigned majority is passed anyway, and a rule that cannot be applied as written cannot be the trust decision. Even where it is applied, what it produces is a blunt, static gate rather than a judgment about whether this particular communication, in this particular context, should be trusted. It cannot tell a first-time caller from a thousandth, a routine request from an anomalous one, or an established account from a compromised one carrying the same valid credential. So the asymmetry the test exposes is not that identity decides nothing. It is that identity can only make a blunt decision, while the relationship layer can decide in context. Remove authentication and a contextual decision still gets made — so authentication was optional. Remove the relationship and all that remains is the blunt gate — so the relationship was load-bearing all along.

The load-bearing test

Remove a component and see if a trust decision still gets made.

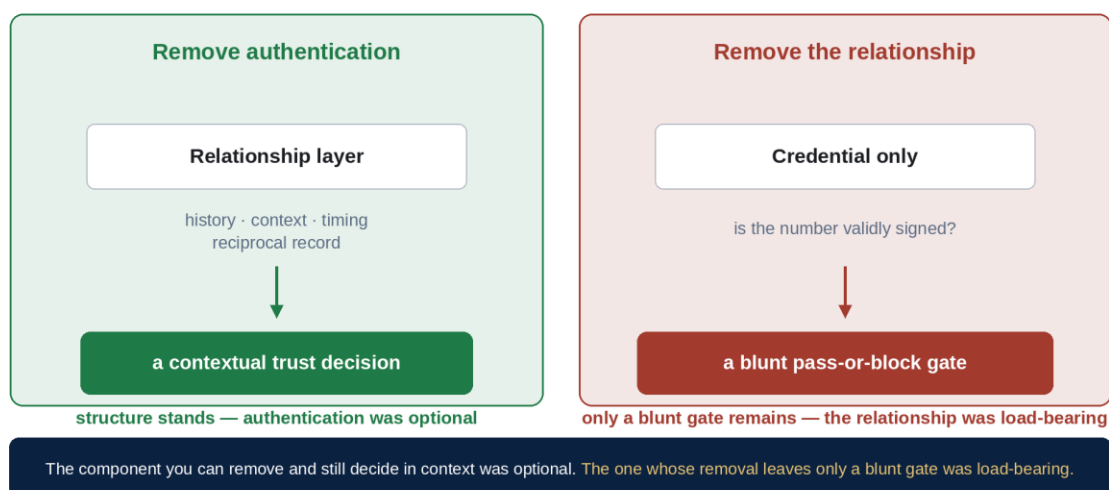


Fig 3-1 · conceptual — no data claimed

The load-bearing test, run in both directions: remove authentication and a contextual trust decision still gets made; remove the relationship and only a blunt pass-or-block gate remains.

Why a credential cannot stand alone: an example

Consider the compromised account. A customer has banked with an institution for many years. The fraud line calls occasionally — always from the same number range, always in business hours, always referencing activity the customer can recognize, never asking for a full credential over the phone. That is the relationship, and both sides hold a version of it. Now suppose the account is taken over and the attacker places a call. Every identity check passes, because nothing about the identity is forged. From the point of view of a pure authentication layer, this call is indistinguishable from the years of genuine ones that came before it — same verified party, same green light.

The relationship layer is positioned to notice what the credential cannot. The credential is valid and the behavior is wrong: a contact at an hour this relationship never uses, a request the real party would never make, a pattern that does not match the reciprocal record the counterparty holds. Identity

verification cannot tell that the account has been taken over — it authenticates the credential, not the behavior, so it has no mechanism to catch a real credential being used wrongly. A relationship layer is not blind in the same way, because it is designed to check whether a communication matches the authorized relationship context and the relationship record; a valid credential acting outside established patterns is designed to surface as a mismatch or anomaly, and the credential alone never confers trust. Overstate the claim, though, and it would repeat the Andersen error in the other direction. The relationship layer is not impossible to trick. It can be tricked, and two later papers in this series are about precisely that. The defensible claim is narrower and holds: the relationship layer is harder to trick, and when it is tricked, it fails less completely — because it watches behavior over time rather than checking a badge once.

Only the relationship can decide alone

Rank the layers on the axis that actually decides the question, which is not importance or order but sufficiency — whether a layer can render a trust decision by itself. Call it the standalone test. It is worth naming what a trust decision is, since everything turns on it. A trust decision is not the same as a disposition: any system can dispose of a call — pass it, block it, ring it through. A trust decision is a judgment about whether this communication fits this relationship, and it cannot be made without evidence about the pair. That is why the credential can be the most valuable input available at first contact, as the previous paper argued, and still not produce one. On that axis the relationship layer is the one layer that passes: it can render a decision without the others, while identity cannot decide without the relationship.

It is worth being precise about what “cannot decide” means, because identity does render a verdict — it simply renders the wrong one where it counts. Among the ten most prolific robocall-signing providers, 97.4% of calls carry A-level attestation, the top grade the framework can issue. The credential is not silent on that traffic and it is not uncertain about it: it certifies, correctly and at the highest grade, the very calls the system exists to stop. The certification is accurate. Treating it as a trust decision is the error. A layer that renders its highest verdict on the worst traffic is not deciding badly. It is answering a different question and having its answer mistaken for a decision. Two properties drive the result. The relationship signal is dynamic, multidimensional, and self-updating, where the credential is static and binary. And the relationship signal is designed to become sufficient once mature — not merely to add confidence at the margin, but to carry the entire decision. The relationship layer leans on identity early and ends superior to it: at cold start the credential is the better signal — the relationship has nothing yet to read; once the relationship is established, it is designed to outweigh the credential and to stand without it. That is a ranking on sufficiency, not a claim about position in a stack, and where it has not yet been proven at production scale it is stated as a design intention rather than an accomplished fact.

That is the conclusion. The rest of this paper asks whether it survives the strongest objections to it: the incentives of the layer being asked to bear the decision, the coverage a perfect credential still cannot reach, and the moment before any relationship exists.

What Each Layer Can Do

Four situations, two layers — the capability matrix behind the load-bearing test.

SITUATION	AUTHENTICATION <i>identity credential — yes/no check</i>	RELATIONSHIP <i>accumulated history — contextual decision</i>	WHAT THIS MEANS
1. COLD START <i>First contact; no history exists yet</i>	STRONGEST MOMENT The best signal going, and nearly the only one. <i>Capability: High</i>	NONE YET No history to read. The first interaction is designed to begin one. <i>Capability: None (by design)</i>	Authentication carries the early decision because there is nothing else.
2. MATURE RELATIONSHIP <i>History exists; many prior interactions</i>	UNCHANGED The same yes/no it gave on call one. <i>Capability: Constant</i>	DESIGNED TO CARRY MOST OF THE DECISION Hours, systems, request shapes, reciprocal record. <i>Capability: High and rising</i>	The relationship layer now carries most of the decision because it has more relevant evidence.
3. COMPROMISED ACCOUNT <i>Credential is real; an attacker is using it</i>	PASSES — THE FAILURE MODE The credential is valid; the identity layer cannot see the misuse. <i>Capability: High — but blind to misuse</i>	DESIGNED TO SURFACE A MISMATCH OR ANOMALY Behavior outside established patterns is designed to surface. <i>Not a detection guarantee</i>	The relationship layer is positioned to notice what the credential cannot. It reduces risk; it does not eliminate it.
4. CAN IT STAND ALONE? <i>Run the removal test in both directions</i>	NO Certifies where a call entered; trust is a question it was never built to answer. <i>Stand-alone: No</i>	THE CLAIM OF THIS PAPER Pull the credential out of the picture; a contextual decision is still designed to be made once history exists. <i>Stand-alone: argued here — by design, with history</i>	Only the relationship layer can, by design, support a trust decision on its own once it has history.

One layer is an input. One layer is designed to be the decision.

Conceptual — capability characterizations in the design register; the compromised-account row states the canonical designed-to-surface claim, not a detection guarantee.

Fig 3-2 · decision capability matrix

Fig P3-2 — What each layer can do, across the four situations this paper examines. It is a map of capability and coverage — what each layer is for, and where each one runs out. It does not by itself prove sufficiency; the table later in this paper does that.

Why the identity layer's incentives are backwards

There is a structural reason not to rest the whole trust decision on the identity layer, and the strongest version of it does not come from us — it comes from the government's own description of the market. Originating providers, the parties who attest to a call's identity at the source, are compensated on the volume of calls they carry, not on the integrity of those calls. The party the industry is asking to vouch for a call is paid to move calls, and paid the same whether the call is honest or dishonest. That is a built-in conflict of interest at exactly the layer being asked to bear the trust decision. It does not make attestation worthless; it means the incentive to attest accurately is structurally weaker than the incentive to attest at all. And the evidence is exactly what that incentive predicts: as the standalone test already showed, the top attestation grade is routinely held by prolific robocallers. The grade is not failing — it certifies where a call entered the network, and it does that accurately. Whether the call should be trusted is a question it was never built to answer. A layer with that incentive structure can be a useful input. It is a poor candidate to be the whole decision.

Perfect authentication still leaves a gap

Grant every other layer its strongest possible form — STIR/SHAKEN attestation signed and surviving end to end, branded caller ID, flawless know-your-customer, no laundering in transit. An independent relationship layer still covers something none of them touch: whether a verified party is behaving consistently with the relationship it claims. That is a gap the others do not close on their best day, because it is not the question they are built to answer. Adding a layer that closes it is therefore additive in coverage. And the gap is not small. It is exactly where the compromised account lives: valid credentials, genuine identity, wrong behavior — a case authentication cannot touch,

however perfect. The size of that share is the subject of a later paper in this series; what matters here is that no amount of authentication, however perfect, can shrink it.

Additive does not mean free. Any additional decision layer can produce false positives: a genuine call slowed or challenged because the pattern looked briefly wrong. That cost is real and has to be managed; a serious deployment is judged partly on how low it keeps the false-positive rate. But it is a tuning problem, not a structural objection: the layer can be tuned toward caution or toward permissiveness, while the gap it closes cannot be reached by hardening anything else at all. So the real question is never whether the relationship layer is perfect — it is whether the coverage it uniquely provides is worth the operational load it adds. For a gap nothing else reaches, it generally is. A defense-in-depth layer whose one real cost is false positives it can be tuned to minimize is difficult to argue against on the merits — which is why this frame, and not any ranking, is the one that holds even for a reader who concedes nothing else.

The cold-start objection, answered

The strongest objection to all of this is cold start: if the relationship layer needs history, what does it do on the very first interaction, before any history exists? Three answers, and they compound. First, cold start is exactly where authentication is most valuable, as the previous paper argued — the layers are complementary, strongest at different moments, not rivals for the same one. Second, at carrier scale the relationship graph does not launch empty. Anchor history already held in the network can be backfilled, so the graph is designed to launch with meaningful breadth rather than starting cold for every pair. That backfill is breadth, not longitudinal depth: it does not manufacture years of history where none exists. What it does is narrow the cold-start gap in practice, so the window in which the relationship layer has little to say is smaller and rarer than the objection assumes. Third, even a truly unknown first contact does not force a blind decision. The interaction can be screened rather than simply allowed or refused — a few questions put to the caller, the exchange visible to the recipient, and the recipient's own decision to block, screen, or trust becomes the first entry in the relationship record. The very first interaction is designed to produce relationship evidence, not merely to happen without it. Even with no authentication layer at all, that disposition still gets made — and, on the definition above, a trust decision is not yet possible, which is precisely why the interaction is screened rather than guessed at.

What Happens to the Trust Decision When a Layer Has Nothing to Read

Both layers are always there. Each answers from its own evidence — and from nothing else. Read down: two inputs, two answers, one conclusion.

	ALL FOUR CALLS — THERE IS NO FIFTH			
	1 — SIGNED NO HISTORY	2 — SIGNED WITH HISTORY	3 — UNSIGNED WITH HISTORY	4 — UNSIGNED NO HISTORY
	UNSIGNED — 51.2% OF ALL CALLS (JUN 2026)			
The call this describes	A signed call — any grade — from a party you have no history with.	A signed call from a party you have history with.	An unsigned call from a party you have history with.	An unsigned call from a party you have no history with.
Is there a credential to read?	YES	YES	NO	NO
Is there any relationship to read?	NO	YES	YES	NO
The credential layer's answer	Verdict: pass. Signed means pass — any grade clears the bar, including the prolific robocaller carrying A-level attestation.	Verdict: pass. Right, but not because it knows anything: it passes for the same reason it passed the robocaller — a valid signature.	Verdict: block. Unsigned means block: the office you have called for years, turned away for a missing signature.	Verdict: block. It reads absence of a signal as evidence against the caller. The rule cannot be applied as written — the unsigned majority is passed anyway.
The relationship layer's answer	Nothing to read — so it asks. No history exists. The call is screened, and the interaction begins the record.	Contextual decision. It reads the history — hours, systems, request shapes — and weighs the credential as one more input.	Contextual decision. It reads the same history, and answers the same way. The missing credential changes nothing.	Nothing to read — so it asks. It will not fake a decision: the screening creates the evidence that was missing, and the recipient decides — accept, block, or keep screening.
Can a trust decision be made?	NO	YES	YES	NO

The relationship is necessary and sufficient. The credential is neither.

A verdict is not a trust decision — and when a layer cannot read, it asks.

Conceptual — structural result in the design register, not measured outcomes. Red marks a verdict that lands on the wrong call. A verdict is the rule's answer, not what the network does: unsigned calls are not in fact blocked. Screening is not a trust decision — it is what happens when none can be made. Signed calls carry one of three attestation grades (A 30.8%, B 4.4%, C 7.8% of all calls); the rule does not distinguish them. Sources: TransNexus (Jun 2026) — see the Evidence companion; how the unsigned 51.2% divides between parties with and without history is not measured.

Fig 3-3 · what each layer answers

Fig P3-3 — What each layer answers, across all four calls that can arrive. Where the capability map above catalogs what each layer is good for, this table asks what follows when a layer has nothing to read: the relationship is necessary and sufficient for a trust decision; the credential is neither.

What this changes in practice

If the load-bearing layer is the relationship rather than the credential, the right place to invest is not another increment of authentication hardening — which, as the coverage and enforcement papers show, runs into structural limits — but the layer that can render a contextual decision and improve with every interaction. It reframes what a trust system is for: not proving who is on the line, which

the AI era has made cheaper to fake and cheaper to verify, but deciding whether this communication fits the relationship it claims. That is a different question, aimed at a different target, and it is the question the papers that follow are built to answer. Standing alone is a stronger claim than eventually outweighing, and this paper treats it accordingly: the structural case is made here; the evidentiary weight rests on the papers and the production record still to come.

House on sand

One objection remains: that a relationship layer built above an unauthenticated network is a house on sand. But run the removal test again: take authentication away, and the relationship layer still decides. A house does not stand when its footings are pulled. The relationship layer does not rest on the credential the way a house rests on its footings; it uses the credential when useful and stands without it when necessary. Optional input, load-bearing layer — the credential can be pulled and the decision still gets made, which is the entire definition of optional. That is why the remaining papers can set the relationship layer aside for stretches and still find the identity approach coming up short on its own terms: short on coverage, where fraud routes around the hardening, and short on enforcement, where even a flawless rule arrives after the call it was meant to stop is already over. The credential matters most at the first hello and matters less with every interaction after it; the relationship is designed to carry the decision alone long before anyone thinks to pull the credential away.

Drafted with AI assistance; all analysis, claims, and conclusions rest with the series editor.

The Permanent Coverage Gap

Why Hardening Call Authentication Cannot Reach the Edges Where Fraud Routes

TRUST IN COMMUNICATIONS · PAPER 4 OF 7 · ICA AI, Inc. · +Trusted Infrastructure · July 2026

A +Trusted Infrastructure series from ICA AI

Series editor: John Perkins, Chief Executive Officer

The previous paper argued the relationship layer can stand alone; this paper sets it aside entirely and grants the identity approach every technical improvement it seeks. The argument is not that attestation failed. It contributed to a measured drop in number spoofing — a win it shares with enforcement, analytics, and labeling. The argument is that hardening attestation cannot reach the places fraud actually routes through, because what keeps it out of those places is economic, not technical. A mandate that lands at origin cannot close a gap that opens in transit and offshore. Coverage thins exactly where the incentive to authenticate is weakest, and the gap that leaves is permanent in a precise, bounded sense: it persists for as long as the economics that create it do, and nothing currently proposed changes those economics.

What authentication got right

Caller authentication — in the United States, STIR/SHAKEN — did real work. Attestation raised the cost of the crudest number spoofing and gave networks an accountability hook they did not have before. This paper does not call that a failure, and the distinction matters: the claim here is that attestation is narrow and insufficient, which is a different and more defensible statement than that it did not work. It worked, within bounds. This paper is about where those bounds are, why they are set by economics rather than engineering, and why that means they cannot be pushed outward by hardening the standard.

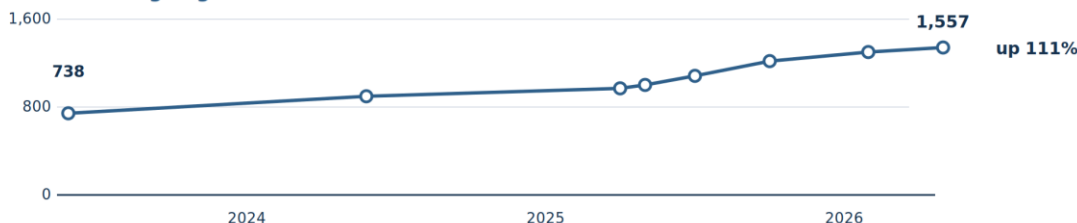
Participation up, coverage still short of half

The measured facts are uncomfortable for the hardening thesis. Even as more providers participate in call authentication, the share of calls arriving fully signed at termination has not kept pace — 48.8 percent as of June 2026, a recovery toward the 49.3 percent peak of October 2024 after a two-year slide, and still short of half. The majority of calls arrive without full-confidence attestation. Participation has risen for five years; coverage has yet to cross the halfway mark — and the recent recovery is the transport gap closing, exactly as the all-IP transition predicts, not the trust gap moving. And over the same years the coverage curve rose and stalled, while reported U.S. fraud losses climbed without pause — from \$5.8 billion in 2021 to a record of roughly \$16 billion in 2025. That divergence is the shape of the whole problem. It resolves into two distinct gaps, and a call is covered only if it clears both — signed at origin, and still carrying that signature when it arrives.

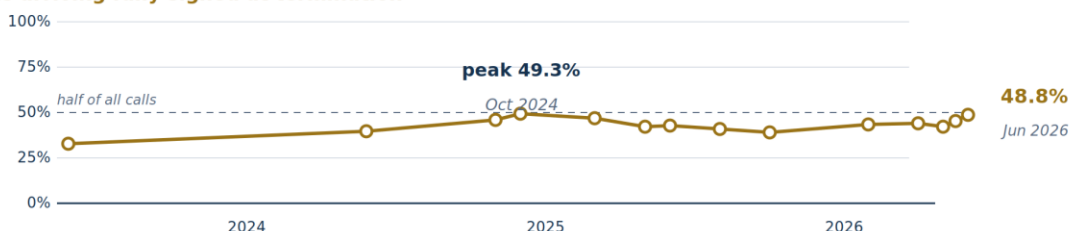
Participation doubled — coverage has yet to cross half

Providers signing calls have more than doubled since mid-2023. The share of calls actually arriving fully signed sits at 48.8 percent — a 20-month high, still short of half after five years of mandate.

Providers seen signing calls



Calls arriving fully signed at termination



Points are TransNexus monthly values; lines connect the measured points.

**Five years of climbing participation — and coverage has yet to cross half.
Participation is not the same as effective coverage.**

Fig 4-1 • TransNexus monthly STIR/SHAKEN statistics, Jun 2023 – Jun 2026

The gap this paper is about, measured. Two lines that were supposed to move together, and did not — which is why participation is not a proxy for the coverage a trust decision would need.

The first is the origin gap: calls that are never signed at all, because the originating segment had neither the incentive nor the capability to sign them. The second is subtler: a call can be correctly signed at origin and still arrive unsigned. As it crosses segments of the network that are not end-to-end IP — older interconnects, certain international handoffs, gateways that translate between protocols — the signature can be stripped or lost in translation: the signature is built to ride on internet-era calls, and an older segment has no place to carry it — so the gateway passes the call along and the proof stays behind. The call that left the originating provider with a valid, top-grade attestation arrives at its destination carrying nothing. This is not hypothetical: of the calls that are signed at origin, only 44.5 percent reached the destination with the signature intact as of February 2026. More than half of the verification is lost somewhere in transit, and a call the system believes it can vouch for arrives as a call it cannot. The figures behind this section, and their sources, are collected in this series' evidence companion, Voice Fraud: The Evidence.

The economics, not the technology

Why does coverage thin where it does? Follow the money: the pattern is not random. The segments where fraud concentrates — under-capitalized, high-churn, often offshore originators — are precisely the segments least able to afford, and least incentivized to deploy, full authentication. Signing calls costs money and effort; a fly-by-night originator that expects to be renamed and reconstituted within the year has no reason to spend either. Hardening the standard does nothing to

that calculus. It raises the bar in the places already over it and leaves the places under it exactly where they were.

Picture a dam built taller across the main channel of a river while the water simply routes around the low banks. The channel is better protected; the total flow downstream is barely reduced, because the flow was never mostly in the main channel. Attestation hardens the main channel — the well-resourced, compliant, largely domestic core of the network. Fraud is in the low banks, and always was. Building the main-channel dam higher is not wasted, but it is aimed at the water that was already contained, and it does not touch the water going around. The analogy has one limit worth naming, and it cuts the wrong way for the hardening thesis: water finds the low bank by accident, while fraud looks for it. An adversary chooses its route. If anything, the picture understates the problem.

Where hardening reaches — and where it does not

Attestation is hardened at the origin. Only one of the three routes a call can take passes through it.

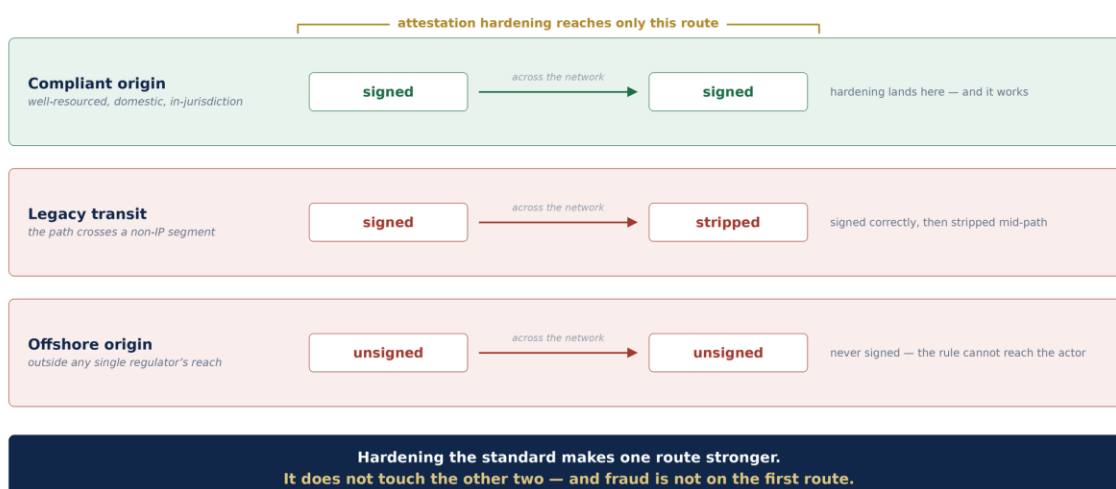


Fig 4-2 · schematic — the three routes are not drawn to scale, and no share is claimed for any of them

Why the gap is structural rather than a drafting error. The mandate lands exactly where compliance is already affordable — and fraud is not on that route.

The conflict at origin

There is a reason the incentive gap persists at origin, and again it comes from the government's own record rather than from us. The originating providers asked to authenticate are, by the regulator's own description of the market, compensated on the volume of calls they carry, not on the integrity of those calls. That is exactly why the top attestation grade is routinely held by prolific robocallers — among the ten most prolific robocall-signing providers, 97.4 percent of calls carry A-level attestation, the highest grade the framework issues. The grade distribution across all traffic tells the same story. Measured against all calls, A-level accounts for 30.8 percent, B for 4.4, and C for 7.8 — 43.0 percent carrying a graded attestation, against the 48.8 percent that arrive signed. When calls get signed at all, they overwhelmingly get signed at the top: of the traffic that carries a grade at all, nearly three-quarters carries the highest one. Authentication is being issued, at the highest grade, to the very actors it was meant to deter, because the party issuing it is paid to carry calls, and the grade it issues

certifies where a call entered the network, not whether it should be trusted. No hardening of the grading standard changes an incentive that runs the other way; a more rigorous grade issued by a party paid to issue grades is a more rigorous stamp on the same incentive.

From the attester's own seat, none of this is irrational — the provider is responding exactly to the incentive it was given. An originating provider paid on volume wants to carry as much traffic as it can and to keep its customers, who are the callers. Assigning a high attestation grade keeps traffic flowing and customers happy; it costs the provider nothing directly when a high-graded call turns out to be fraudulent, because the grade certifies origin, not conduct, and the consequences of the conduct land downstream on the recipient. So the rational move, for a provider optimizing the incentive as written, is to attest generously. The grade drifts upward not because anyone is cheating but because the system pays for volume and asks for integrity, and when those two pull apart, the party in the middle follows the one that pays. This is the deep reason the top grade ends up on the worst traffic: it is the predictable output of the incentive, not an aberration in it, and a standard that grades more strictly without repricing the incentive simply produces stricter grades pointed at the same calls.

Mandates cannot fix affordability

The industry's answer is regulation — the FCC's proposed know-your-customer and know-your-upstream-provider requirements, detailed in this series' companion on the proposed rules. Grant that they help; they raise the floor and close some of the crudest origin gaps. They still do not touch the affordability problem. A mandate can require authentication, but it cannot make authentication economical for an under-capitalized offshore originator, and it cannot reach a foreign segment outside its jurisdiction. And these requirements are, as of this writing, proposals — origin-level, years from full effect, with continuous know-your-customer still a proposal in a rulemaking rather than adopted law. A rule that lands at origin cannot close a gap that opens in transit, where signatures are stripped, or offshore, where the rule does not reach. The mandate and the gap are, structurally, in different places.

A gap with a passport

The affordability gap has a geographic twin, and the two reinforce each other. A large share of the fraud that reaches domestic phones originates in, or is laundered through, segments of the network that sit outside the reach of any single national regulator. A rule written in one jurisdiction binds the providers it can name; it does not bind a foreign originator, and it does not bind the intermediate carrier in a third country that hands the traffic along. The call arrives domestically, but the point where it should have been authenticated sat in a place the rule could never touch. This is not a loophole to be closed by tighter drafting — it is a boundary of sovereignty. Authentication mandates are national; the fraud economy is not, and a gap that opens outside the jurisdiction cannot be closed by a rule inside it, however well written. The border gap and the affordability gap fail the same way: the rule lands in one place, and the party who would have to comply is in another.

What hardening can buy, and what it cannot

The case for simply doing more hardening conflates two different things: what a stronger standard can buy, and what it cannot. Hardening can buy three real things: a higher cost for the crudest impersonation, a firmer accountability hook inside the compliant core of the network, and a cleaner signal on the calls that are signed and stay signed. Those are worth having. What hardening cannot buy is the larger remainder of the problem. It cannot make signing economical for an originator that expects to disappear and reconstitute under a new name. It cannot reach an actor outside the jurisdiction. And it cannot keep a signature intact across a segment that strips it. No increment of rigor moves an item from the second list to the first, because the items on the second list are not failures of rigor — they are failures of reach, economics, and sovereignty, and a better standard does not address them. This is why the debate about how much more to harden, however earnest, is aimed slightly past the problem: it is optimizing the part that already works while the part that does not sits outside its reach.

Why the gap persists

The gap is durable — stronger than temporary, short of eternal — because its cause is structural, not incidental. Costs could fall; attestation could reach further; some of this could change. What does not change easily is the shape of the problem. Fraud concentrates at those same edge segments — the pattern that enforcement records and the robocall-signer data both point to — and they are precisely where the economics of authentication break down. And between origin and destination, the legacy segments that strip signatures in transit are receding — carriers are retiring copper and TDM against real schedules — but they are not disappearing uniformly or all at once, and the international and edge segments where fraud is most likely to route are the last to convert. That transit gap shrinks with the all-IP transition. It is the smaller of the two gaps, and the transitional one.

The larger gap is not about transport at all. Even on a network that is fully IP end to end, where every signature survives and every eligible call is signed, authentication would still certify only where a call entered the network and under what claimed authority — not whether the communication should reach the person on the other end. Grant STIR/SHAKEN its best possible future — complete IP coverage, universal signing, no transit loss — and the trust decision is still unmade. The transport gap shrinks. The trust gap does not.

No amount of hardening at the well-resourced center changes the economics at the under-resourced edge. Permanent, then, is shorthand — precise shorthand: the transport gap fades with the transition, but the trust gap beneath it persists for exactly as long as authentication is asked to answer a question it was never built to answer. That leaves one more question worth asking on the identity approach's own terms. Suppose the proposed rules worked perfectly anyway — flawless rollout, full compliance. Would even that be enough? The next paper grants exactly that assumption.

Figures cited in this paper

Share of calls arriving fully signed at termination: 48.8 percent (June 2026), a 20-month high, recovering toward the 49.3 percent peak (October 2024); the majority still arrive without full-confidence attestation. TransNexus.

Of calls signed at origin, the share arriving with the signature intact: 44.5 percent (February 2026). Grade split at termination, as a share of all calls: A 30.8 percent, B 4.4 percent, C 7.8 percent — 43.0 percent graded; among the top-10 robocall signers, 97.4 percent of calls at A-level (June 2026). Reported U.S. fraud losses: \$5.8 billion (2021) to roughly \$16 billion (2025), FTC Consumer Sentinel.

Figures as of this writing; sources and fuller data in the companion, Voice Fraud: The Evidence.

Drafted with AI assistance; all analysis, claims, and conclusions rest with the series editor.

Even If the Proposed FCC Rules Work Perfectly, Will They Be Enough?

Why Know-Your-Customer and Authentication Mandates Cannot Be the Sole Arbiter of Trust

TRUST IN COMMUNICATIONS · PAPER 5 OF 7 · ICA AI, Inc. · +Trusted Infrastructure · July 2026

A +Trusted Infrastructure series from ICA AI

Series editor: John Perkins, Chief Executive Officer

This paper grants the proposed know-your-customer and know-your-upstream rules the most generous assumption available: a flawless rollout and full compliance. The argument does not rely on the relationship layer, and makes no claim about any product. The argument is entirely about the rules on their own terms. Even granting perfect compliance at origin, the enforcement remedy behind these rules is a slow sequence ending in a sanction so severe it is never used quickly — which means a bad actor can operate in the lag. Prevention leaks at the edges by the rules' own scope, and enforcement arrives too late to patch it. That is why compliance, even flawless, cannot be the sole arbiter of which calls get through.

Granting the best case

Assume everything the rules' strongest advocate would ask for. Every originating provider knows its customers. Every upstream provider is vetted. Attestation is applied correctly and consistently across the compliant network. This paper does not argue that outcome is impossible or that the rules are misguided — it stipulates the ideal and asks a narrower question: with compliance working exactly as designed, is the framework enough to serve as the arbiter of trust?

What perfect means here needs defining, because otherwise the argument looks like it wants things both ways — granting perfect prevention and then relying on evasion. It does not. Perfect compliance with the rules as written is not the same as perfect prevention. The rules are origin-level, and the previous paper showed that origin-level coverage cannot reach the edges where fraud routes — the under-resourced originator, the foreign segment, the transit hop that strips a signature. So even a flawless rollout leaves a prevention gap, and it leaves it not through any failure to comply but by the rules' own scope. Grant perfect compliance, then. Evasion still occurs, by scope rather than by breach, and the real question becomes what the framework does about the traffic that gets through anyway.

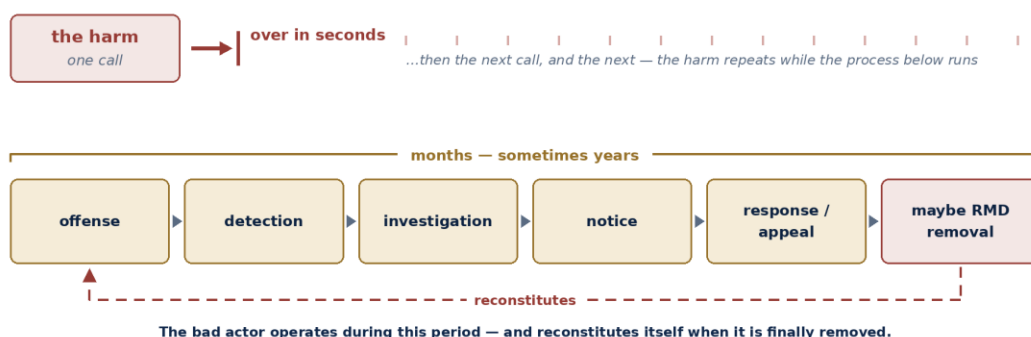
The remedy is a sequence, not a switch

When a bad actor slips through, the rules do not stop them at the moment of the offense. Enforcement is a sequence, and the length of the sequence is the point. Walk it stage by stage. First the violation has to be detected, which itself takes time, since the harm usually surfaces before its source is identified. Then it is investigated, to establish who is responsible and whether the conduct meets the standard for action. Then the provider is noticed and given an opportunity to respond, and to appeal — due process that a fair system owes even to bad actors. And only then, at the far end and only sometimes, comes removal from the Robocall Mitigation Database, the step that effectively ends a

company's ability to originate traffic. Each stage takes time because fairness requires it. The framework is built to prevent and to punish, but the punishment arrives at the end of a process measured in weeks, months, and sometimes years — while the harm is measured in the length of a phone call.

Seconds to strike, months to remove

The harm is one call; the remedy is a long, fair, and therefore slow sequence.



A fair remedy is necessarily a slow one.

The severe sanction (RMD removal) is real but never used quickly — it arrives a lap behind the traffic.

Fig 5-1 · conceptual — no data claimed

The offense and the remedy run on different clocks. Nothing in the sequence below is a failure of diligence — every stage of it is there because fairness requires it.

The first stage is usually the longest and the least visible. Detection is not instantaneous, because fraud announces itself through its effects, not its source. The harm surfaces first — a wave of complaints, a spike in a particular scam, a pattern of losses — and only then does the work of tracing it back to a responsible originator begin. By the time the source is identified with enough confidence to act, the operation has often been running for a long stretch, and the damage that triggered the detection is already done. Everything downstream of detection — investigation, notice, appeal, removal — then adds its own time on top. So the lag is not one delay but a stack of them, and the first and largest is the one that runs before anyone even knows where to point.

The severe sanction is never used quickly — for good reason

Removal from the database is the framework's ultimate sanction, and the record shows how deliberately it is used. The first mass removals came in August 2025: an initial 185 providers, then the certifications of more than 1,200 more three weeks later; another 1,203 followed in January 2026 — the first mass sweeps in the database's history. None of it was quick. The compliance deadline the August cohort missed fell in February 2024; the show-cause order reached 2,411 providers in December 2024; removal landed in August 2025 — eighteen months from missed deadline to

disconnection, and a removed provider may refile only with the agency's approval. Set that against the scale of the problem: U.S. consumers received 52.5 billion robocalls in 2025 alone, 29.6 billion of them scam calls, while reported fraud losses reached a record of roughly \$16 billion. It is important to read that record correctly. The pace is not slow because the agency is failing. It is slow because due process and proportionality are real constraints that a fair regulator observes, and because even a mass removal comes only after warnings, cure periods, and a formal show-cause process. The restraint is reasonable. An agency that pulled that sanction casually, on thin evidence, would be acting lawlessly, and the same caution that makes the process fair is what makes it slow. The slowness is not a bug that better administration would fix; it is the cost of doing enforcement fairly, and it is a cost worth paying — which is exactly why the remedy cannot also be the real-time arbiter of a live call.

The enforcement sequence — from missed deadline to disconnection

The rule was enforced. The record shows it was never used quickly: eighteen months from the deadline to the first removals.

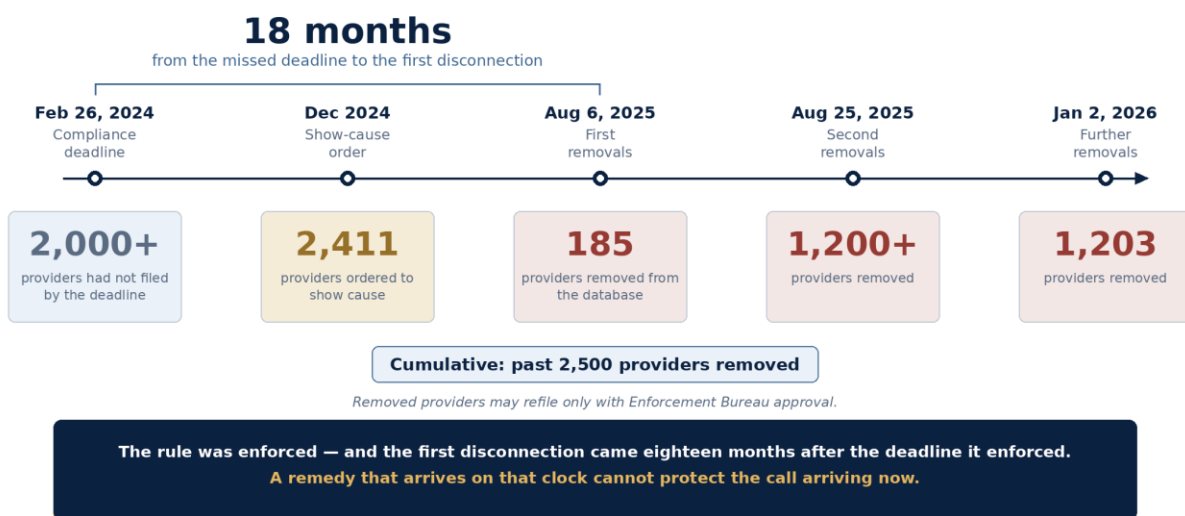


Fig 5-2 · FCC Enforcement Bureau orders, Aug 2025 – Jan 2026

The sequence is not hypothetical — it has run, at scale, on the public record. What the record shows is the clock, not the will.

The lag is the opportunity

A bad actor does not need the enforcement process to fail. They need it to be slow, and it is slow by design. So they operate inside the lag — the interval between the offense and the sanction — and they treat that interval as the business model rather than a risk to it. Spin up an originating entity, run traffic through the gap while the clock on detection and investigation and appeal runs its course, and when removal finally lands, the entity has already earned out. The mass-fraud economics of the space make it cheap to reconstitute under a new name and begin the clock again. Meanwhile the policy response to a genuinely new evasion is often another rulemaking, which runs on the timescale of months and years while the traffic runs on the timescale of seconds. Prevention is a rule; fraud is a race — and the remedy is always running a lap behind, catching the runner who has already finished.

Why removal does not end the traffic

Even when the sequence runs to its end and the sanction lands, the traffic does not necessarily stop, because the entity that was removed is not the same thing as the operation behind it. Removal ends a company's registration; it does not end the people, the capital, or the know-how that ran it. In a market where standing up a new originating entity is cheap and fast, the rational response to removal is not to exit the business but to reconstitute it — a new name, a new registration, the same operation, and the enforcement clock reset to zero. The regulator, having spent months of process to remove one entity, then faces the same operation again under a different label and must begin the sequence anew from detection. This is not a failure of the removal; the removal did exactly what it was designed to do. It is a mismatch of speed and cost. The sanction is expensive and slow to apply, and the thing it removes is cheap and fast to replace, and a remedy that costs the enforcer more than it costs the offender to evade cannot, on its own, win a volume game. The math of the contest favors the side that can reset the clock, and that is the side the rules are trying to stop.

A referee who can only review the goal

The structural picture is a referee who is not on the field during the play but reviews it afterward from the replay booth — and who, by rule, may overturn the call only in the clearest cases, and rarely does. Every part of that is fair. The referee is impartial; the review is careful; the standard for reversal is appropriately high, because reversing a result is a serious act. And the goal still counts, far more often than not, because by the time the review concludes the play is long over and the score already stands. Nothing about this describes a broken referee. It describes a referee built to review rather than to intervene — and a system that decides trust only in the replay booth cannot decide, in the moment the call is live, whether this particular call should be trusted. The competence of the review is not the issue. Its timing is. The series never asks whether the offender is eventually punished; it asks whether the person receiving the call should trust it before answering.

Prevention leaks, enforcement lags

This paper and the one before it are two halves of a single argument. The previous paper showed that prevention leaks at the edges, because the economics of authentication break down exactly where fraud routes. This paper shows that enforcement lags because a fair remedy is, by necessity, a slow one. Those are the only two things a compliance regime can do — prevent and punish — and each has a structural gap the other cannot close. Prevention cannot reach the edge; enforcement cannot arrive in time. Between them they leave a space, and the space is not incidental — it is exactly the space in which a live trust decision has to be made: this call, right now, before it connects and before any enforcement process could possibly run.

What a live decision actually needs

What falls into that space can be named. A live trust decision needs something present at the moment of the call, specific to the parties on it, and able to weigh behavior rather than only origin. Prevention offers a pre-call credential that says nothing about conduct; enforcement offers an after-the-fact judgment that arrives too late to inform the call it concerns. Neither operates at the moment and place where the decision is actually made. This is not an argument that the rules should be weaker or that compliance is pointless — it is a real and valuable improvement, and granting it its perfect

day is the fairest test available. It is an argument that a framework which can only prevent at origin and punish after the fact has, by its own structure, left the live decision unmade. Granting the rules their perfect day does not close that gap. It clarifies where it is.

Drafted with AI assistance; all analysis, claims, and conclusions rest with the series editor.

The Proposed FCC Rules

Know-Your-Customer, Enforcement, and the Gaps That Remain

TRUST IN COMMUNICATIONS · COMPANION VOLUME · ICA AI, Inc. · +Trusted Infrastructure · July 2026

A +Trusted Infrastructure series from ICA AI

Series editor: John Perkins, Chief Executive Officer

Drawn from the FCC public record (CG Dockets 17-59 / 02-278, WC Docket 17-97; FCC 26-27, FCC 26-32). Figures current as of July 2026; the docket posture is live and may change.

This companion to the Trust in Communications series explains the rules the industry is treating as its answer to voice fraud — what they require, who has to carry them out, how they are enforced, and, above all, when they actually take effect. Drawn entirely from the public record, it reaches a plain conclusion. The requirements that would genuinely tighten the network are still proposals in an open comment cycle. They run at the origin, where a provider vets who may put calls on the network, and never at the relationship between the two parties on a call. And the path from proposal to a rule taking effect runs for years. A solution that isn't enforceable for years, and that structurally cannot reach the offshore edge, leaves the coverage gap open for a long time. The rules raise the floor — which is worth doing — but by their design they do not decide whether a given communication should be trusted.

Why they need something new

Start with the question the proposals themselves answer: if the existing rules worked, why write new ones? The network has not been unregulated. For years, originating providers have been required to know their customers, to authenticate calls through the STIR/SHAKEN caller-ID framework, to support one-day traceback, and to block categories of clearly illegal traffic. That framework is real, and it did real work. And yet, across those same years, reported fraud losses kept climbing and the share of calls people least want kept rising. The existence of a new, more demanding proposal is itself the admission: the rules in place did not close the gaps they were built to close. This is not to say the old rules failed outright — they narrowed some kinds of abuse — but the abuse that mattered most routed around them. Everything that follows is the FCC's attempt to reach further. The question this companion asks is how much further it can actually reach.

Adopted law versus proposal

The single most important thing to hold onto is the difference between what is law today and what is merely proposed, because the industry's confidence borrows from the second while the network only runs on the first. What is adopted law today is a baseline: originating providers already must take effective measures to know their customers and their upstream partners, every provider in the call path must file and stand behind a robocall-mitigation plan, and the FCC enforces this with real tools — removal from the mitigation database (which cuts a provider off from the U.S. network), traceback, cease-and-desist orders, and monetary settlements such as the Lingo Telecom AI-voice case. That baseline is genuine.

What is not yet law is everything that would turn that general duty into a detailed checklist. Today's rule tells providers to know their customers but leaves the how to them; the proposals would replace that with an exact set of required steps — specific information to collect, verify, retain, and re-check. All of it lives in Further Notices of Proposed Rulemaking adopted for comment in spring 2026, with no final rules and no effective dates. When someone says the FCC “is requiring” these steps, they are describing a proposal in the present tense. That single tense shift is where the grand-solution story quietly overstates the record.

The Know-Your-Customer proposal (FCC 26-27)

The centerpiece is the Know-Your-Customer proposal, adopted for comment on April 30, 2026. It would convert today's general duty into a concrete framework for originating providers — the phone companies that place calls onto the network, from a national carrier to a small VoIP or wholesale operator. Before such a provider may let a customer — the business or bulk caller buying the right to originate calls — make calls, it would have to collect at least the customer's name, physical address, government-issued identification number, and an alternate telephone number; verify that information; retain the records for four years; and re-verify when red flags arise. High-volume callers would face heavier checks, and the FCC declined to let providers rely on vetting that someone else already performed.

Two features of that list deserve a closer look, because the gap between what a rule says and what it can do lives in the details. The first is the government-issued identification number. Collecting a number is not the same as verifying it against the authority that issued it, and the proposal leans toward collection. A number is a string of digits: a determined actor can supply a stolen real one, a fabricated one that passes format checks, or a synthetic identity assembled from both. Unless the provider checks it against the issuing government's records — which is hard to do at scale and often impossible across borders — the requirement stops casual fraudsters but not determined ones, and a foreign fraudster hands a U.S. provider an identity it has no practical way to confirm.

The second is the promise to re-verify when “unusual activity” or “a change in traffic patterns” arises. That sounds proactive, but it raises an obvious question: who is watching, and how? The proposal names the red flags without saying who monitors for them, against what baseline, or how quickly — so the re-verification is only as good as a monitoring capability the rule assumes but does not define, fund, or standardize. The instinct behind it is revealing: reaching for changes in behavior over time is a move toward watching conduct rather than checking a credential once. But a provider watching a customer's aggregate traffic sees volume and destinations in bulk; it has no view into whether any single call fits the relationship between that caller and the person receiving it. The instinct points past a one-time check; the vantage point stops at the origin.

The proposal also asks whether prepaid and postpaid customers should be treated differently — and that this is still an open question is itself telling. Postpaid service is billed after the fact, which leaves a credit check and a paper trail; prepaid is paid up front with minimal identity, easily bought in bulk and discarded, which is exactly why a large share of fraud originates there. Leaving prepaid's treatment unresolved means the highest-risk, lowest-identity channel does not yet have a settled

standard; if it lands lighter, the rule writes its own bypass. The deeper point is that prepaid is a lawful, mass-market product used by tens of millions, and it is not going to be banned. Applied to prepaid, the effect of KYC is to attach a verified identity to an account that today carries little or none — not to eliminate the product, but to remove its anonymity. Whether that can be done reliably, at scale, and across borders against fabricated identities is the open question — and for the same reasons the identification check is weak, it is a question the origin-level approach struggles to answer.

The companion upstream proposal (FCC 26-32)

Riding alongside it is a Know-Your-Upstream-Provider proposal, taken up in May 2026. A call often passes through several carriers before it reaches a phone; the carrier that hands it to the next one is “upstream” of the carrier receiving it, and the originating carrier is furthest upstream of all. This proposal would require each provider to vet not just its own direct customer but the upstream provider that handed it the traffic — collecting business, ownership, and financial information, checking the upstream provider’s mitigation filings, and verifying what it claims — pushing accountability back up the chain. It would also tighten how the certificates that let a provider participate in call authentication are issued and revoked, and, importantly, codify the A, B, and C attestation levels with specific criteria. The proposal would turn those grades from a convention into a rule.

What the attestation levels certify — and how common each is

Most calls carry no grade at all. Even the top grade — the most common of the graded calls — certifies only the number, not that the call is honest.

	Authenticated customer relationship	Verified right to use the number	Share of all calls
No attestation the largest group — no readable grade at all	×	×	≈57%
A Full attestation knows the customer and confirms the number	✓	✓	≈31%
B Partial attestation knows the customer, not the number	✓	×	≈4%
C Gateway attestation cannot authenticate the source at all	×	×	≈8%

100% of all calls

“No attestation” means no readable grade — which is not the same as unsigned. It includes calls that were signed at origin and lost the grade in transit: 51.2% of calls arrive unsigned, and the graded calls (A 30.8 / B 4.4 / C 7.8) sum to 43.0%, leaving 57% with nothing to read.

Every grade certifies where a call entered the network — not whether it is wanted or honest.
Even A, the top grade and the most common, can ride a fraudulent call.

Fig R-1 · share of all calls, TransNexus, Jun 2026

The distinction the grade hides: A-level attests the number was authorized, not that the caller is honest — which is why a prolific robocaller can hold A. Share of all calls; TransNexus, Jun 2026.

Who has to deliver it — and who is out of reach

The obligation falls on originating voice service providers on the U.S. network, and the proposed scope is broad, not limited to the largest carriers: wireline, mobile, interconnected VoIP, SIP-trunking and wholesale originators, calling-platform providers, and mobile virtual network operators — with particular attention to prepaid, bulk-SIM, and high-volume outbound accounts, where the crudest abuse concentrates. So the answer to who operationalizes it is: every originating provider touching U.S. traffic, each vetting its own customers and upstream partners as business entities.

What no version of the rule reaches is just as important. The rules vet identity — who a caller is, and whether they may put traffic on the network. They never reach the relationship — whether a given communication fits the history between the two specific parties on it. They do not reach a foreign originator or a transit carrier in a third country, which sit outside the FCC’s jurisdiction. And they do nothing about signatures stripped mid-path: a call signed correctly at its source can still arrive stripped, because the break happens in the middle of the path while the mandate sits at its start. The rule lands at the origin, in one jurisdiction, on identity; the fraud routes around all three of those limits.

Who has to deliver it — and who is out of reach

The scope at the origin is broad. The origin is one point on a path the rule does not otherwise touch.

ORIGIN — IN SCOPE		every originating provider touching U.S. traffic
●	Wireline carriers	THE RULE LANDS HERE and only here
●	Mobile carriers (CMRS)	
●	Interconnected VoIP	
●	SIP-trunking and wholesale originators	
●	Calling-platform providers (CPaaS)	
●	Mobile virtual network operators (MVNOs)	
Each vets its own customers and its upstream partners as business entities — with particular attention to prepaid, bulk-SIM, and high-volume outbound accounts.		
THE REST OF THE PATH — OUT OF REACH		NO RULE REACHES HERE and this is where fraud routes
×	Transit carriers, mid-path where the signature is stripped — the mandate sits at the start of the path, the break is in the middle	
×	Foreign originators outside the FCC’s jurisdiction; no U.S. rule reaches them	
×	Third-country transit carriers outside the FCC’s jurisdiction	
×	The relationship itself no version of the rule asks whether a call fits the history between the two parties on it	

The scope is broad across the origin — and the origin is one point on the path.

Fig R-2 · who the rule reaches, and who it does not

Six kinds of provider must vet their customers and upstream partners at the origin — but the rest of the path carries no such duty, and that untouched stretch is where the fraud routes.

What the rules reach — and the holes they leave

Even fully adopted and enforced, the obligation sits at one layer, in one jurisdiction, on a multi-year clock.

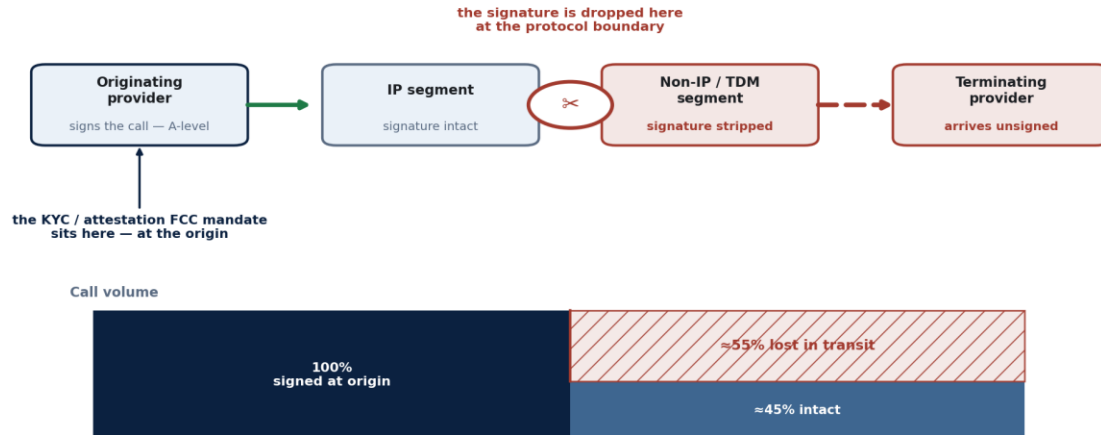


Fig R-3 · what the proposed rules reach, and what they leave open

The holes are structural, not gaps to be patched later: each one follows from where the obligation is placed, so tightening compliance narrows none of them.

How a signed call arrives unsigned

A call can be signed correctly at the origin and still arrive stripped. The break is mid-path; the FCC mandate sits at the origin.



The signature is created and verified at the origin.
A rule that lands there cannot protect it across a break in the middle of the path.

Fig R-4 · conceptual — mid-path signature loss

About 55% of signed calls lose the signature at the mid-path break — an origin-level mandate cannot reach the segment where the loss happens.

How it is policed, and what failure costs

Policing runs through the same machinery that exists today, strengthened. A provider that fails is exposed to per-call forfeitures, to traceback that walks a bad call back to its origin, to cease-and-desist and blocking directives, and ultimately to removal from the mitigation database — the step that cuts a provider off from originating onto the U.S. network. The proposal sets that forfeiture at \$2,500 per illegal call, tying the penalty to the volume of bad traffic. That final sanction is not a rarely-used relic; lately it has been used aggressively. In August 2025 the Enforcement Bureau removed 185 providers, then over 1,200 more three weeks later; further removals followed, and in January 2026 it cut off another 1,203. To size that against the whole: roughly 9,200 providers are registered in the database. The December 2024 show-cause cohort alone — about 2,400 flagged for deficient filings, some 1,400 ultimately removed — came to roughly fifteen percent of that universe, and the January tranche pushed cumulative removals past 2,500, disconnected in a matter of months.

That volume is real. But two things put it in proportion. First, this was the first time the tool had been swung at scale: the removals cleared a backlog of deficient filings that had built up for years under looser rules, after a 2024 order tightened the standard and directed those 2,400 providers to cure or be cut. A first-time catch-up is not a repeatable engine — a backlog can only be cleared once. Second, the removals were overwhelmingly for paperwork — inaccurate or incomplete certifications — not the slow, case-by-case adjudication of a provider proven to be originating fraudulent traffic, which still runs through detection, investigation, notice, and appeal before any sanction lands.

What comes next is the harder problem, and it is different from the cleanup. Small originating providers serve only a sliver of real subscribers, but they carry a disproportionate share of illegal traffic — which is exactly why they are the target, and why the removals hit real fraud volume rather than a rounding error. But they are cheap to start and quick to replace. Remove one and its fraud customers route to the next; the operator reconstitutes under a fresh filing while the clock resets. And the reconstitution need not even be sequential. Nothing in the framework stops one operator from registering several entities at once and holding the spares in reserve — in which case removal costs no time at all, because the successor is already filed and already registered. The enforcement sequence assumes a gap between being cut off and coming back; parallel registration closes that gap before it opens. The database can be swept clean of who is registered today. It cannot stop who registers tomorrow. Whether enforcement can keep pace with that turnover, rather than clear a one-time backlog, is the open question the removal numbers do not answer.

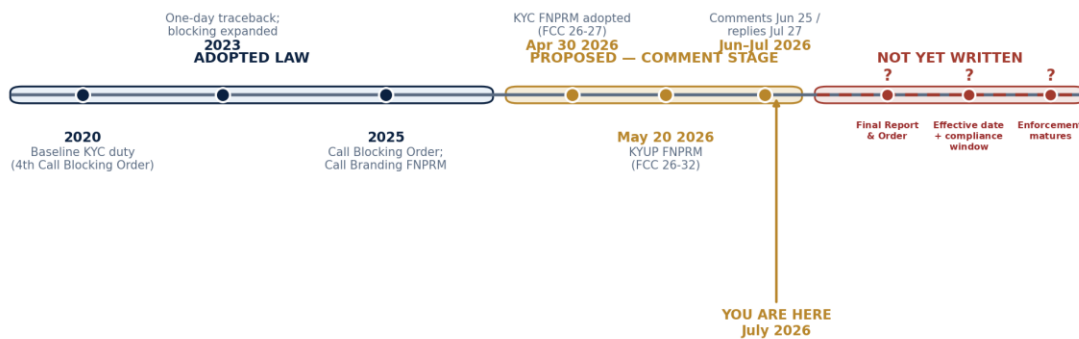
The chronology

Here is the timing, which is the heart of the matter. The baseline duty dates to 2020; one-day traceback and expanded blocking came in 2023; further blocking and call-branding items came in 2025. The two proposals that matter here were adopted for comment only in spring 2026 — the Know-Your-Customer item on April 30, with comments due June 25 and reply comments July 27, and the upstream item in May. As of this writing the record is still in its comment window. Nothing after that is scheduled: the FCC must still write and adopt final rules, set effective dates, and phase in compliance — each of which takes months to over a year, and effective dates in this area have a

documented habit of slipping once set. The realistic path from today's proposals to rules that actually bind and are actually enforced is measured in years.

The chronology of the “grand solution”

Adopted baseline → today's proposals → the years before anything takes effect



Every requirement that would actually tighten the network is still a proposal. The path from proposal to a rule taking effect runs for years — and effective dates, historically, slip further still.

Fig R-5 · FCC rulemaking chronology

The stretch between today's proposals and a rule that actually binds is unscheduled and long — and that lag is the interval a bad actor operates inside, the same gap Paper 5 turns on.

What the record adds up to

Set the findings side by side, drawn only from the record:

- They are proposals, not law — and years away from taking effect.
- They act at the origin, vetting who may put calls on the network — never at the relationship between the two parties on a call.
- They cannot reach the offshore originators and transit segments where much of the fraud routes, which sit outside any single national rule.
- Their central identity check can be satisfied with an identity that was never verified at its source.

Each is a way of saying the same thing: the rules raise the floor, and the floor is worth raising, but a floor is not a decision about whether a given call should be trusted.

That is not a case against the rules. It is a description of what they leave unsolved. The problem the proposals are chasing — deciding, in the moment, whether a communication is what it claims to be — remains unsolved when the last proposal is adopted, because it sits at a layer these rules do not operate on. The current path can raise the floor but, by its own structure, cannot be enough, and the problem is still open. That opening is what invites a different approach — one aimed not at who may originate a call, but at whether the call fits the relationship it claims. ICA AI is pursuing exactly that approach, and says so plainly here; this companion makes its case on the public record alone and

leaves the reader to judge whether the gap it documents is real. If it is, closing it is work that still has to be done — attempted, and proven — by someone. That much the record makes clear.

Drafted with AI assistance; all analysis, claims, and conclusions rest with the series editor.

Why a Relationship Cannot Be Faked at Scale

Bidirectionality, Economics, and the Limits of Impersonation

TRUST IN COMMUNICATIONS · PAPER 6 OF 7 · ICA AI, Inc. · +Trusted Infrastructure · July 2026

A +Trusted Infrastructure series from ICA AI

Series editor: John Perkins, Chief Executive Officer

A fair critic will ask the obvious question: if trust rests on a relationship, can an attacker not simply fake the relationship? This paper takes that objection seriously. A single relationship can be mimicked, for a time, by a determined attacker — this paper does not claim otherwise. The defensible claim is narrower: a relationship cannot be faked at the scale that makes fraud profitable. The load-bearing reason is structural — bidirectionality, the fact that the record lives on both sides and across many independent counterparties — and it is the reason that holds even as AI drives the cost of imitation toward zero. Scale-economics reinforces it but is not, by itself, the barrier. The barrier is structural and belongs to the relationship itself — it holds wherever two parties have a history, whoever is keeping the record, and whether or not anyone builds a layer to read it.

Two questions, not one — and only scale matters

The objection collapses two different questions that need to be kept apart, and most of the confusion about relationship-based trust lives in the gap between them. Can a relationship be faked? Yes — with enough effort against a single target, an attacker can study a pattern and mimic it for a while. Can a relationship be faked at scale? That is the question that matters, because fraud is a volume business. An attack that works once but does not repeat is not a fraud operation; it is a stunt. Fraud pays because it is cheap per attempt and runs across millions of attempts, so the only forgery that threatens a trust system is one that scales. This paper concedes the first question without reservation and contests the second, and the whole argument is about why the answer to the second is no.

The distinction tracks how the two sides actually operate. A defender — the party fraud is aimed at — does not need every individual relationship to be unforgeable. They need forgery to be a losing business at the scale that makes fraud profitable. An attacker, conversely, does not win by faking a single relationship; they win only if the forgery repeats cheaply across a large population of targets. So the whole contest comes down to scale, and a claim about a single relationship — forgeable or not — settles nothing. The right question is whether the property that resists forgery holds up when the attacker tries to run it a million times, and that is a question about structure, not about any one relationship's defenses.

A ladder of forgeries

Impersonation is a ladder: each rung is harder than the last by a different kind of difficulty. The first rung is forging the number — spoofing caller ID — which is cheap and thoroughly solved for the

attacker; this is the problem the identity era was built to address. The second rung is forging one side of a relationship — studying how a trusted party behaves and reproducing its texture: the right hour, the right references, the right rhythm. This is harder, and AI is making it cheaper, which is precisely the critic's strongest point. But there is a third rung, and it is a different problem, not a harder one: faking both sides of the relationship, consistently, across every counterparty that holds its own memory of the dealings. The first two rungs are effort problems, and effort is exactly what better tools reduce. The third is not an effort problem at all, and that difference is the subject of this paper.

The attack gets cheaper — until it hits the wall

*Two steps are effort problems, and AI is driving the cost of effort down.
The third is a structural problem.*

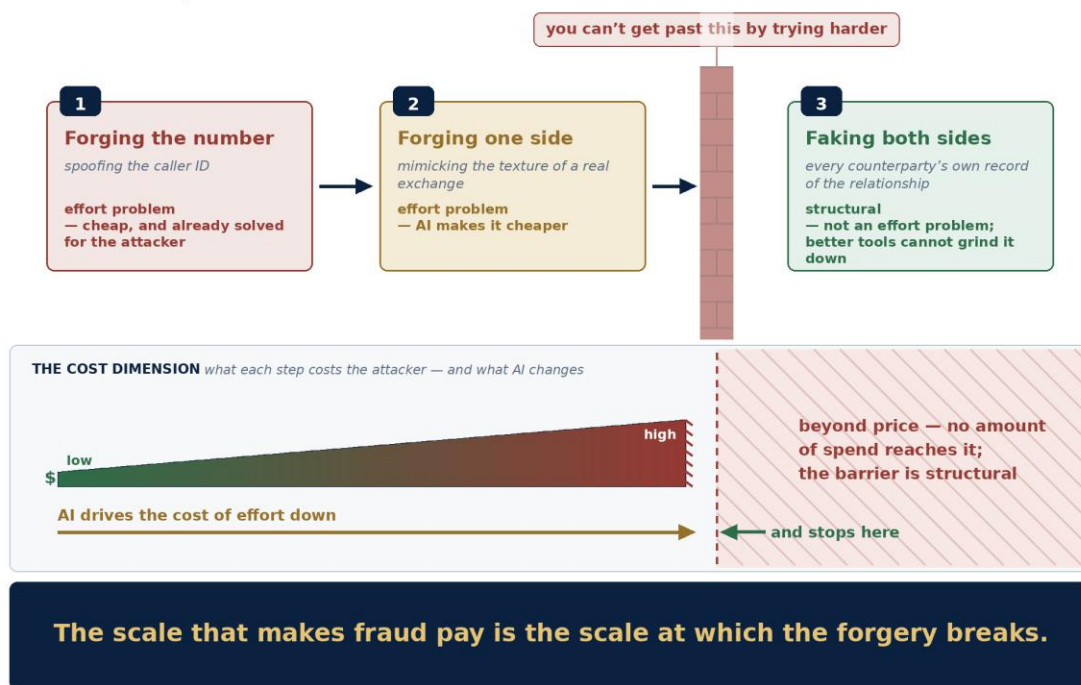


Fig 6-1 · conceptual — no data claimed

The ladder, and where it stops being a ladder. The first two rungs are bought with effort; the third is not for sale at any price.

Bidirectionality: a relationship lives on both sides

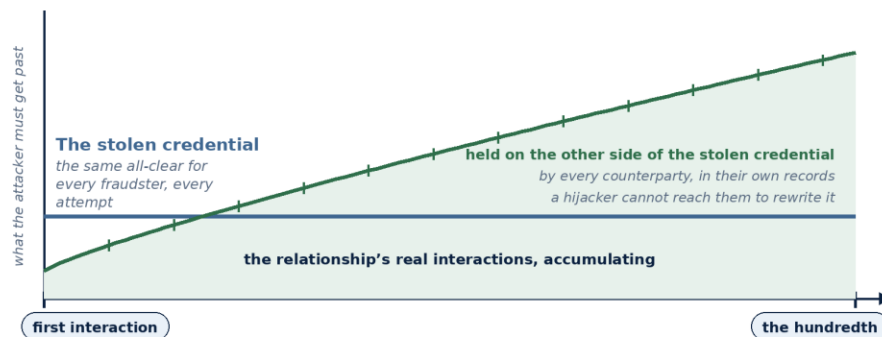
The third rung is hard because a relationship is not a credential the caller carries; it is a record both parties hold, and the two records have to agree. An attacker impersonating a trusted party can control their own side of the story completely — but not the counterparty's. To fake the relationship convincingly, they would need to write themselves into the other party's record too, and then into the records of everyone that party deals with, because those cross-references are designed to form part of what the pattern is checked against.

Picture a business with five hundred established relationships, or a person with fifty. That identity is not a badge either one carries; it is scattered across all of those separately held records — five hundred or fifty, the number does not matter — each maintained by a different counterparty, each

reflecting a slightly different history. An attacker forging that business or that person controls none of those records. Each one is another independent account the forgery has to survive — records the attacker does not see, does not have practical access to, and does not control. Forging one side is imitation; forging both sides, across a web of counterparties, is a different and far harder problem, because the evidence does not live in one place where a single forgery can reach it. Better tools cannot grind this barrier down: it is a matter of where the records live, and they live in hands the attacker will never hold.

The same stolen credential — against a record that keeps growing

*A hundred fraudsters can buy the same valid credential; it opens the same door every time.
What they must match is the relationship record — and every real interaction adds to it.*



A stolen credential is the same lie every time.

A relationship record is a hundred facts the attacker never saw.

Fig 6-2 · conceptual — the axes carry no scale

What the attacker buys, against what the attacker must match. One of these is fixed at purchase; the other keeps growing in hands he will never hold.

The record the attacker cannot reach

Why are the counterparties' records out of reach? Not because they are encrypted or well-defended, though they may be. It is that they are held by other parties entirely, maintained independently, and updated by those parties' own real interactions. When a business genuinely calls one of its five hundred counterparties, that counterparty's record updates from its own side — its own observation of the interaction, on its own system. The attacker impersonating the business has no way to produce that update, because producing it would require being the counterparty, not the business. This is the difference between forging a document and forging a fact that many independent witnesses separately recorded. You can forge the document you hold; you cannot reach into five hundred other parties' memories and revise what each of them independently saw. A relationship, distributed across the records of everyone who participates in it, is closer to the second than the first — which is why controlling your own side, however perfectly, does not forge the whole. And those records are

not isolated from one another: each is shaped by its keeper's own dealings with parties the forger never touched, so what a forgery has to match is not a set of separate targets to be taken one at a time, but a shape produced by behavior the attacker has no part in.

The spoofed-context attack, taken seriously

The strongest form of the objection is the spoofed-context attack: the attacker does not just spoof the number but mimics the context — calls at the usual hour, references the expected account, matches the rhythm of a real relationship. Even done well, the mimicry is built from the outside. Reconnaissance harvests what the relationship has shown — the hours, the references, the rhythm. A forgery assembled from the visible surface is thinner than the relationship it imitates, and that is part of why it works only for a time. This is a real attack, and against a single target with enough reconnaissance it can work for a time. But notice exactly what it costs. The attacker has to learn one relationship deeply enough to reproduce its texture, and having done so, has faked exactly one — and only its visible side, the side that faces the target. And none of that work can be reused. The next target is a different pattern, held by different counterparties, with a different history; it requires the same hand-built effort from scratch. The attack that succeeds is the attack that does not scale, and the attack that would scale is the one that cannot succeed, because it would require reaching records the attacker does not control.

It is worth being exact about what the relationship layer reads, because the spoofed-context attack is aimed at something it does not look at. The layer does not read content. It does not listen to the call, parse the message, or interpret what was said. What it reads is the shape of the exchange — who contacted whom, when, through which channel, in which direction, at what cadence, and whether the counterparty's own record agrees. Metadata and math: structured signals, and the arithmetic over them. An attacker who perfectly reproduces the substance of a conversation has reproduced the one thing the layer never consults — and still has to reproduce a pattern that is held on both sides of a relationship they were never part of.

This is a security property and a privacy property at the same time. A layer that never reads content cannot leak content, and cannot be quietly repurposed to surveil it. The same thing that makes a relationship hard to forge — that it is shape, not content — is what keeps the substance of the conversation out of the record entirely.

Why it does not scale — including with AI

There is an easy version of this paper's argument, and it goes like this: fraud pays because it is cheap and repeatable — send the same lure to a million numbers and profit on the fraction that respond. Faking a relationship is nothing like that. It is hand-built, one target at a time: expensive, slow, unrepeatable. So — the easy argument concludes — relationship forgery will never scale. All of that is true today. And all of it is fragile, and a sharp critic will say exactly why: it is an argument about human effort, and AI erodes human effort. The same evidence base this series draws on shows AI collapsing the cost of hand-built voice attacks, with deepfake fraud up more than thirteen-fold.

So the easy argument is not where the barrier actually lives, and it should not be asked to carry the case. The barrier is bidirectionality, and AI does not touch it. A model can generate a plausible pattern for one side of a relationship; it has no practical way to write the attacker into the independent records that hundreds of counterparties each hold and check against. AI sharpens the point rather than blunting it. A model can spin up a complete-looking web of relationships for nothing — but complete-looking to whom? The fraudster has no knowledge of the truth it is meant to match, and the records it will be checked against sit in other parties' hands. Fabricate the web and it is free, and wrong in ways the attacker cannot see. Research the web instead — scrape the public trail, cross-reference the connections — and the cost returns, and the harvest is still only the surface the relationship chose to show. Either way, the attacker works without an answer key.

Forging the number is cheap. Forging one side of the story is getting cheaper. Forging both sides, consistently, across every counterparty's separately held record, is not an effort problem that better tools reduce — it is a structural requirement that cannot be met, because the evidence the attacker would need to control does not live anywhere the attacker can reach. That is why the claim is that a relationship cannot be faked at the scale that makes fraud profitable, not that it can never be faked. Economics tells you the attack does not pay against many targets; bidirectionality tells you it does not work against them, and the second is the claim AI cannot erode.

Designed to surface the mismatch, not to catch the attacker

The honest register for how this shows up in practice matters, because it is easy to slide from a structural argument into an operational promise the architecture does not make. The claim is not that the system catches every forgery or blocks it outright. It is that behavioral divergence from an established, bidirectional pattern is designed to surface as a mismatch or anomaly — a signal that gets weighed, not a verdict that gets pronounced. A patient, well-resourced attacker may narrow the divergence for a while against a single relationship, and against that one relationship they may succeed. What they cannot do economically, or structurally, is narrow it across the many relationships that fraud at scale requires, each held by different counterparties who each hold their own record. This paper establishes only why the structural property exists; how a system surfaces and weighs it is outside the scope of this paper. The barrier is bidirectionality. Fraud pays by repeating cheaply. This forgery starts from scratch every time, so it never gets cheap enough to become profitable at scale.

Drafted with AI assistance; all analysis, claims, and conclusions rest with the series editor.

The Compromised Account

When a Real Account Is Taken Over and the Credential Is Genuine

TRUST IN COMMUNICATIONS · PAPER 7 OF 7 · ICA AI, Inc. · +Trusted Infrastructure · July 2026

A +Trusted Infrastructure series from ICA AI

Series editor: John Perkins, Chief Executive Officer

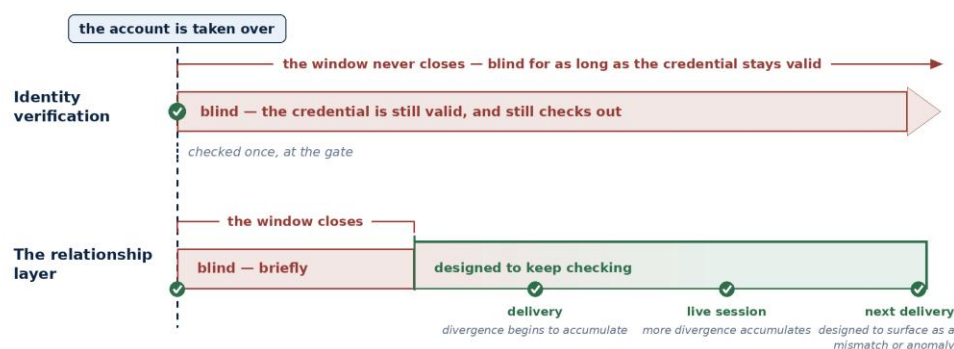
Every trust approach has a weakest edge; this paper is about the relationship approach's. When a real account is taken over, there is an opening window in which the compromise has not yet shown itself — and the relationship layer does not claim to close that window instantly. The case is a comparison: identity is blind to account compromise for as long as the credential stays valid, which is indefinitely; the relationship layer is blind only until behavior diverges enough to surface, and it is designed to keep watching. A one-time gate versus a continuous check.

The stolen account: the attack that passes every check

Take the case that is genuinely hard for any behavioral approach: a real, established account is compromised, and the attacker operates it using entirely valid credentials. Every identity check passes, because nothing about the identity is fake — the account is real, the credential is real, and only the party behind it has changed. Paper 3 used this case to prove the relationship layer can render a decision alone; this paper stays with it, because it is also where the relationship approach is tested hardest. And for some interval after the takeover, the attacker's behavior may not yet have diverged far enough from the established pattern to raise a flag. That interval is the opening window, and this paper is not going to pretend it does not exist or that a clever enough design makes it disappear. Some fraud is fast; a window of even a short duration is a real exposure. A trust system that claimed to close this window entirely would be making exactly the kind of over-promise this series argues against.

Blind briefly — or blind forever

The account is real and the credential is valid. What differs is how long each layer stays blind after it is taken over.



Identity verification is blind for as long as the credential stays valid — indefinitely.
The relationship layer's blind window is designed to close as divergence accumulates.

Fig 7-1 · conceptual — no data claimed

The hardest case for this series, drawn: both layers see a real account and a valid credential. Only one of them has anything left to check.

What this layer does not claim

The relationship layer does not claim to detect a takeover the instant it happens, and it does not claim that continuous checking is the same as instantaneous certainty. There is a period after compromise during which a careful attacker, operating a real account and staying close to its established behavior, looks enough like the account holder to pass — that period is the window, and it is the attacker's window of opportunity. The value of the relationship approach has to be measured against the honest alternative, not against an imaginary perfect one — and the case does not require that the window never opens. It requires that the window closes, and closes soon enough to matter. That is the comparison with the alternative, where the window never closes at all. It would be simple to claim the system detects takeovers, full stop. That claim would be false at the edges. The truthful version is longer and less impressive: the system is designed to shorten the interval in which a takeover stays invisible, and how short depends on how quickly the attacker's behavior gives itself away.

Blind briefly versus blind forever

Here is the comparison that matters, and it is the whole paper in one contrast. Identity verification is blind to account compromise for as long as the credential remains valid — which is to say, indefinitely, because a taken-over account presents a real credential and a point-check has no mechanism to notice that the party behind it has changed. For a pure identity approach, the window never closes on its own; it closes only when something outside the system intervenes. The relationship layer is blind only briefly. Because it is designed to check behavior continuously against the established pattern rather than once at the gate, the compromise has a limited time to stay invisible: as the attacker acts, divergence from the pattern is designed to accumulate and surface as an anomaly. The difference is not fast detection versus slow detection. It is a window that closes versus a window that does not — a one-time gate that checks once and then trusts forever, against a continuous check that keeps asking whether the behavior still fits.

Checking more than once

What does continuous actually mean here? A gate checks at one moment — the credential is presented, verified, and the party is trusted from then on. A continuous check does not stop at the gate; it re-evaluates whether the communication still matches the authorized relationship context at later points, including at delivery and during a live session, rather than resting on a decision made once at the start. The credential answers a question asked at the door; the continuous check keeps asking the same question after the party is inside. The check reads the shape of the exchange, not its substance — the same metadata and math Paper 6 describes. A hijacker who sounds exactly right has matched the one thing the layer does not consult. This is why a valid credential never confers standing trust here: the check that mattered at the door is not the only check, and a party that passed the door but then acts outside the established pattern is designed to surface as an anomaly at one of the later points rather than to ride the entry decision indefinitely. None of this makes the surfacing instantaneous. It makes the checking ongoing, which is precisely the property a one-time gate structurally lacks — not faster checking, but checking that does not stop.

Why divergence accumulates

Why does the window close at all? The whole comparison rests on the answer. A hijacker running a fraud operation cannot both use a real account and stay perfectly inside its established behavior indefinitely, for a simple reason: staying perfectly inside the established behavior earns them nothing. An attacker who wants only to watch — espionage, surveillance — is a different problem, and not the one this series is about; fraud has to act. To profit from the compromise, they eventually have to do something the real relationship would not — reach a party it would not reach, at a time it would not, in a shape the relationship has never taken. The everyday version is familiar to anyone with a phone: the taken-over account that suddenly blasts a message to everyone in the contact list. The real account holder never messaged four hundred people at once, and never at that hour. To the counterparties who each hold their own history with that account, the blast is not one anomaly — it is four hundred divergences at the same moment, one in every record it touched. That deviation is the entire point of the takeover; a compromise that never deviates from normal behavior extracts no value. And the response is itself evidence: what each counterparty does with a divergent contact — block, screen, or trust — becomes the next entry in the record the attacker would have had to predict. And each deviation is a divergence from the pattern the counterparties hold. A continuous check is designed to weigh each such divergence rather than wave it through on the strength of a valid credential. A one-time gate cannot do this at all: once it has checked the credential, it has nothing left to check, and the account is trusted for as long as the credential lasts. A continuous check has the rest of the relationship's life to notice that something has changed — and the attacker's need to act is what gives it something to notice.

Where the value is, and where it is not

This protection is not uniform — where it helps most depends on the speed of the attack. Against an instantaneous drain — where value moves in the first moments after takeover, before any behavior has had a chance to diverge — a continuous check has the least time to work, and this paper does not pretend it closes that case cleanly. Against slower relational fraud — where the attacker must build toward the payoff over multiple interactions, each an opportunity to diverge — the continuous check earns far more of its value, because the attack gives it time and material to surface the anomaly. The honest scope, then, is not that the relationship layer solves account takeover. It is that it converts a permanent blindness into a closing window, and that the window closes fastest exactly where the attacker is forced to act over time. The benefit is real. So are its limits.

Honesty is the point

The window is real, and anyone who works in this domain knows it. The relationship approach does not close the window the moment it opens. It is designed to make the window short and closing rather than permanent and open — and that is the whole of the claim. What the divergence is surfaced from is not a mystery: the shape of the exchange, weighed against the record the counterparty holds. What a system then does with that signal — the thresholds, the response — is implementation, and the comparison does not depend on it: one approach checks once and stops; the other keeps checking.

That contrast is where this series ends, because it is the choice the industry now faces. Seven papers have asked one question — whether verifying identity is enough to establish trust — and the answer,

tested from every direction this series could find, is that it is not. Trust lives in the relationship. The infrastructure that reads it is the next layer to build.

The Trust in Communications series — in reading order

Seven papers, two companions that hold the evidence and the rules, and a concluding epilogue.



Read in order, or read the one that answers the question you came with.

Fig 7-2 · series reading order

Where this paper sits: the last of the seven. The companions hold the evidence and the rules; the epilogue follows.

Drafted with AI assistance; all analysis, claims, and conclusions rest with the series editor.

Epilogue

What Changed, What Remains True, and Where the Work Goes Next

TRUST IN COMMUNICATIONS · EPILOGUE · ICA AI, Inc. · +Trusted Infrastructure · July 2026

A +Trusted Infrastructure series from ICA AI

Series editor: John Perkins, Chief Executive Officer

This epilogue introduces no new argument and no new evidence. Seven papers and two companions made the case; what is left is to take stock: what the argument changes, what it leaves standing, where it stops, and what it leaves for others to take up. A reader who has followed a long argument is owed that accounting. We would rather say plainly what is established than let the case overreach on our behalf.

What changed

The series proposes a different engineering model of trust. It does not argue that identity is unimportant; it argues that identity and trust are different engineering objects, and that treating them as the same has constrained how the industry asks its central question.

In practice, in product, and often in policy, the two have been treated as interchangeable: verify the caller, the reasoning goes, and you have settled whether the call can be trusted. But identity is an attribute of a party, and trust is a property of a relationship between two parties. They are held in different places, they are established by different means, and they fail in different ways. That is the change the series is asking for — not that identity matters less, but that it is a different object from trust and cannot be made to stand in for it.

The argument reached that position by elimination rather than accumulation. It asked one question — whether verifying identity is enough to establish trust — and each paper removed a way of answering yes:

- ❑ **Origin is not trust.** What authentication establishes is a fact about the call — that it was signed, and how strongly — not a fact about the party behind it. Trust is a property of the relationship — something a single signed call cannot establish.
- ❑ **The relationship layer can stand on its own.** It can carry a decision by itself, tested against its hardest case — the compromised account.
- ❑ **Hardening the standard does not close the gap.** Coverage divides into two gaps, and only the smaller one is transitional. Signatures stripped in transit recover as the network converts to all-IP — which is what the 2026 recovery in coverage is. The larger gap is untouched at any coverage level: authentication reaches where the economics of signing make it worthwhile, and fraud routes precisely where it doesn't.

The last three removals were the hardest, each taken on the opposing position's strongest ground:

- ❑ **Even a flawless rulebook leaves the trust decision unmade.** Grant the Proposed FCC Rules full rollout and compliance, and rules that can only prevent and punish still never answer the

only question that matters when the phone actually rings — should this call be trusted? Prevention works only in advance, and enforcement only after the fact.

- ❑ **A relationship resists forgery at the scale that makes fraud profitable.** The record is held on both sides, and across every counterparty that keeps its own account of it; an attacker controls only their own side — and AI cannot forge records it was never given.
- ❑ **Even a stolen credential doesn't beat a layer that keeps checking.** A one-time gate is blind as long as the credential stays valid, which is indefinitely; a layer that keeps checking is blind only until the attacker acts — and the attacker has to act to get any value from it.

What is left standing, after all seven papers, is that the decision requires something present at the moment of the call, specific to the parties on it, and able to weigh behavior rather than only origin.

What remains true

The series spent seven papers on what these tools cannot do. That is not the same as saying they do nothing. Plainly: each of them still does real work.

- ❑ **Authentication** answers a question worth answering, and where it is deployed it has measurably reduced spoofing — a real achievement of the last decade.
- ❑ **Regulation** helps — the Proposed FCC Rules are a real improvement at what they set out to do.
- ❑ **Enforcement** matters, and its slowness is not a defect to be engineered away — it is the cost of doing enforcement fairly, and a cost worth paying.
- ❑ **Identity** has a role the series never disputed: a trust decision has to attach its history to a party, and identity is how a history is indexed.

Indexing the history is not carrying the decision. When the series says authentication is not necessary, it means for carrying the decision, not for filing the history — the relationship record can recognize a counterparty even when no credential arrives with the call. Identity files the history; it does not render the verdict.

None of that is in question. The series denies only that these tools are enough — not that they are worth having. Each does its job. But none of them decides, when the phone rings, whether this particular call should be trusted — and that decision is the one the industry still has to make.

Where the argument stops

Four things the series does not claim:

- ❑ **It is an argument about structure, not a system.** Nothing in the series depends on any implementation, and nothing in it should be read as describing one. Where a paper says the relationship layer, it means the property — the thing that is true whether or not anyone ever builds it — not a product.
- ❑ **The claim is that divergence surfaces, not that it is detected.** Behavioral divergence from an established, bidirectional pattern is designed to surface as a mismatch or anomaly. It does not claim detection, and it does not claim that surfacing is the same as certainty. A signal that

gets weighed is not a verdict that gets pronounced, and staying honest about the claim means never letting the first pass for the second.

- **The forgery claim is only about scale.** A single relationship can be mimicked by a determined attacker with enough reconnaissance, and the mimicry can work for a time. But fraud is a volume business, and a forgery that does not repeat is not a fraud operation.
- **The window on a compromised account closes — just not instantly.** That it closes at all is the property a one-time gate structurally lacks, and the series claims nothing more.

Where the work goes from here

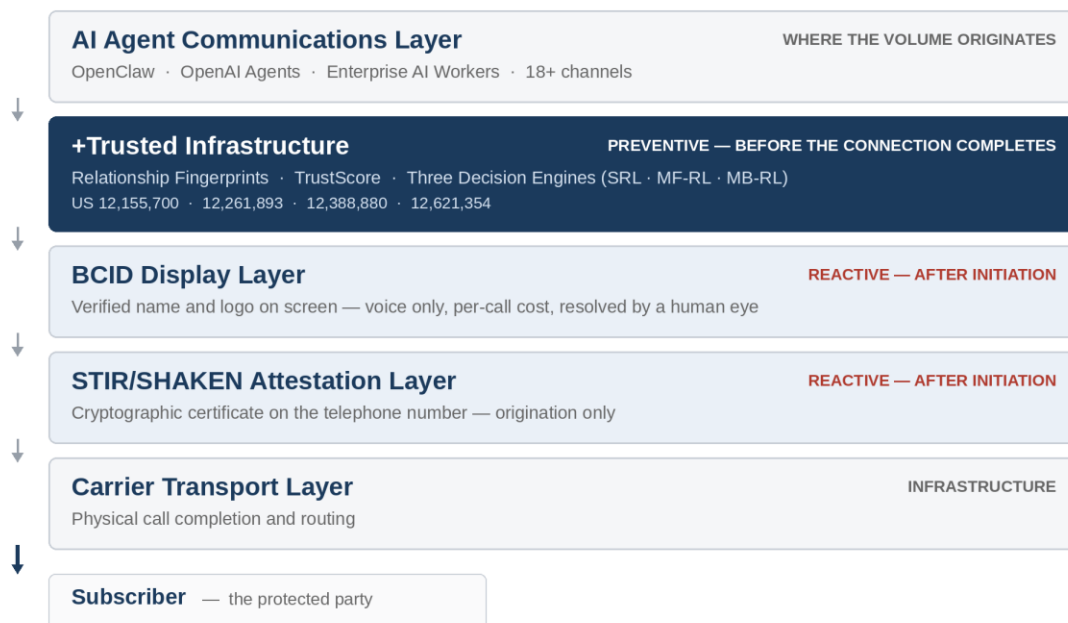
If the series is right that trust is relational, there is real work still to do. The following are research and engineering questions that need to be answered:

- **How should a relationship be represented?** The series argues that relationships are load-bearing without settling what counts as an interaction of record, how much history is required before a relationship signal can carry weight, or how a relationship that has gone quiet should be treated. Until these are settled, there is no common object to build on or to measure.
- **Where should the evidence live, and what privacy must it guarantee?** The barrier in Paper 6 depends on records reflecting what each party independently holds — and wherever that evidence lives, the privacy properties it must guarantee are a real and open question.
- **What happens before a relationship exists?** Every relationship has a first interaction, and at that moment there is nothing to read. The series' answer is to screen that first contact and let its outcome become the first entry in the record. What remains open is narrower: how much history counts as enough before a relationship signal can carry weight.
- **How much does behavior have to break pattern before it counts — and what does getting it wrong cost?** Surfacing a divergence is not the same as acting on one. Thresholds, the cost of a false positive against a false negative, and what a system may do with a weighed signal are live questions — and they are the ones on which operator adoption will actually turn.
- **How would anyone measure whether any of this works?** Today the industry measures authentication — signing rates, attestation levels, database compliance — and those process metrics are the ones it relies on. Nobody measures whether the person receiving the call was protected. Roughly 48.8% of calls were signed at termination as of June 2026 — a twenty-month high, and still short of half after five years of the mandate — and of the calls signed at origin, fewer than half arrive with the signature intact; reported fraud losses hit a record of roughly \$16 billion in 2025. The first two numbers describe a mechanism, the third an outcome — and no one has an accepted way to connect them. No approach to trust, this one included, can be judged against a benchmark that does not exist.

The question changes

If the arguments in this series are substantially correct, then the industry's central research question changes. It is no longer how do we strengthen identity. It is how do we responsibly operationalize relationship.

That work is the domain of +Trusted Infrastructure — ICA AI's layer built for precisely the characteristic that authentication does not have: the relationship.



The communications stack: only the preventive layer decides, before the connection completes, whether a communication should reach the subscriber.

This series does not claim to have answered that question. It claims to have changed it. The answer belongs to those who would build it — to the operators who would carry it, the researchers who would measure it, and the standards bodies who would have to agree on what is being measured. Seven papers asked whether verifying identity is enough to establish trust, and the answer, tested from every direction this series could find, is that it is not. That finding does not close the matter. The work it opens is a research program, and it belongs to the people who will build the answer, not to this series.

Colophon

Trust in Communications, monograph edition — the complete series in reading order. A +Trusted Infrastructure series from ICA AI, compiled July 2026. The contents and reading order appear at the front of this volume. The individual papers and companions are published at ICATrusted.ai. Drafted with AI assistance; all analysis, claims, and conclusions rest with the series editor. © 2026 ICA AI, Inc.