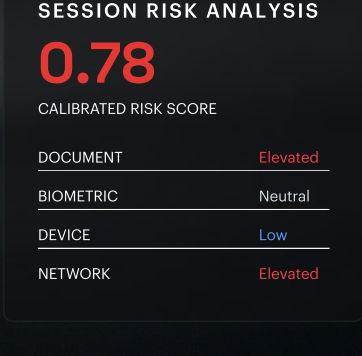


Fraud Hub: Multi-Signal Risk Scoring

AI-powered AML compliance and transaction monitoring that detects sanctions hits, catches suspicious patterns, and keeps your compliance team audit-ready — all through one integration.



The Problem Fraud Hub Solves

Most verification platforms evaluate fraud as a sequence of independent checks. A document module returns pass or fail. A liveness module does the same. A device check adds its verdict. A rules engine combines the flags. Each module compresses its analysis into a single binary output before the next stage sees it.

That compression destroys forensic detail. The camera stream entropy, the document compression artefact, the device driver signature, the IP-to-document geolocation delta: gone before any cross-domain analysis is possible. Gartner predicts that by 2026, AI-generated deepfakes on face biometrics may cause up to 30% of enterprises to lose confidence in standalone verification (Gartner, 2024).



Signal Loss Through Compression

When a device fingerprinting service returns "high risk," the specific browser anomalies, geolocation variance, and timing patterns that distinguish a VPN user from a device spoofer are discarded. The scoring model never sees them.



No Cross-Domain Correlation

A document that passes in isolation should be flagged if the device is emulated, the IP routes through a residential proxy, and the facial biometrics show micro-inconsistencies. Siloed tools cannot correlate across domains.



Grey-Zone Fraud Exploits the Gap

The most expensive fraud sits between obvious pass and obvious fail. Synthetic identities and coordinated fraud rings produce signals that individually look borderline but collectively reveal manipulation.



Vendor Sprawl Multiplies Risk

Managing separate contracts for device fingerprinting, liveness, document verification, and network intelligence creates integration complexity, support fragmentation, and multiple potential points of failure.

"Each signal individually looks acceptable. Analysed together, before any of them are compressed into a binary output, the attack is unambiguous."

Evaluate whether your current verification stack detects coordinated, multi-vector fraud.

Request a Demo

Talk to a Fraud Expert

What It Means for Your Team

Fraud Hub addresses different priorities depending on the role evaluating it. Production deployment with a major financial services client validated both sides of the accuracy equation simultaneously.



SECURITY & FRAUD TEAMS

Detect Grey-Zone Fraud at Scale

ML analysis across four signal domains detected 70% more fraudulent documents than binary rule-chaining in production deployment. Calibrated scoring provides auditable, explainable decisioning backed by forensic signal decomposition.



COMPLIANCE & RISK

Audit-Ready Score Explainability

Every Risk Score decomposes into contributing signals across all four detection domains. No black-box outputs. Supports regulatory explainability requirements under the EU AI Act, PSD3, and AML6D with a transparent, calibrated probability rather than arbitrary thresholds.



PRODUCT & ENGINEERING

Zero Integration Overhead

Fraud Hub activates within the existing Shufti verification flow with zero integration overhead. Production deployment showed a 57% reduction in false rejections of legitimate users, directly improving conversion rates.



FINANCE & PROCUREMENT

Consolidate Vendors. Quantify ROI.

Four natively owned detection domains in a single contract, replacing the need to source device fingerprinting, liveness detection, document verification, and network intelligence from separate providers. Each falsely rejected user is lost revenue. Reducing false rejections by 57% has a directly quantifiable revenue impact.

Fraud Hub: Four Domains. Raw Signals.

Shufti natively owns all four fraud detection domains. The ML engine ingests raw forensic signals from each domain simultaneously, before any signal is compressed into a binary output. When these domains are outsourced to third-party providers, the underlying forensic data does not survive the handoff. This is the structural gap Fraud Hub closes.

DOMAIN	LOST WHEN OUTSOURCED	ANALYSED BY FRAUD HUB
Injection Attack Detection	Camera driver signatures, payload timestamp anomalies, and hardware ID inconsistencies — discarded at the third-party handoff before any cross-domain pattern can form.	Camera stream entropy · Virtual camera driver signature · Hardware device ID consistency · Payload header authenticity · Metadata timestamp sequence
Device, IP & Network	Session-specific device fingerprints, cross-client fraud linkages, and the relationship between network path and document geolocation. Absent from a third-party IP risk score.	Device fingerprint hash · OS build integrity · VPN/Tor exit node match · Datacentre/residential/mobile classification · IP-to-document geolocation delta · Session velocity
Document Presentation Attack	Raw compression artefact patterns, metadata tool signatures, and pixel-level edge anomalies that correlate with injection signals. Absent from a verified/rejected verdict.	Compression quantisation (QTable) deviation · Metadata modification tool signature · MRZ check digit validation · Template deviation score · Security feature presence · Pixel-level manipulation artefacts
Facial Liveness & Deepfake	Micro-pattern characteristics of specific deepfake generation tools, which correlate with injection attack signatures from the same session. Cannot be cross-referenced from a pre-scored liveness output.	Liveness challenge response quality · Ambient light reflection consistency · Frequency-domain artefact signatures · Facial landmark geometry · Skin texture micro-pattern · Eyes-open-and-tracking confirmation

A Calibrated Risk Score

Every session produces a single calibrated Risk Score where the numerical value directly represents fraud probability. A score of 0.7 means a 70% likelihood of fraud, not an arbitrary threshold or opaque tier. This calibration allows precise threshold-setting for different risk appetites. In the verification report, the fraud analyst sees the calibrated score alongside a per-domain signal breakdown: which signals fired, at what intensity, and how they correlated across domains. The score is not a number in isolation; it arrives with its own forensic explanation.

1

Parallel Signal Collection
All four domains collect signals simultaneously during the standard verification flow. No additional SDK. No extra API calls. No user-facing friction.

2

Raw Signal Ingestion
The ML model receives granular forensic data from each domain before binary compression. Cross-domain correlations remain intact.

3

Dynamic Signal Weighting
Signal weights shift based on context. A VPN flag in isolation is treated differently from a VPN flag that co-occurs with a document geolocation mismatch and a virtual camera signature.

4

Explainable Output
Triggered signals surface alongside the Risk Score in the verification report. Every score decomposes into contributing factors across all four domains for audit and compliance.

Assess how Fraud Hub performs against your current fraud detection stack.

Book a Technical Walkthrough

Where It Applies

High-risk decision points where standalone verification leaves the most dangerous gaps.

Account Opening
Synthetic identities and coordinated fraud rings targeting volume onboarding flows with fabricated documents.

Tier Upgrades & Limit Increases
Mule accounts seeking higher transaction thresholds using stolen or manipulated credentials.

Withdrawals & Transfers
Account takeover and stolen identity cash-out attempts at the point of highest financial exposure.

Account Recovery
Identity impersonation during credential reset, a common attack vector for third-party fraud.

Enhanced Due Diligence
Step-up identity proofing where standard KYC is insufficient for regulatory or risk purposes.

Ongoing Transaction Review
Fraud patterns emerging post-onboarding from accounts that passed initial verification.

The Architectural Difference

Most fraud scoring products operate on signals that have already been compressed. They receive binary outputs from separate modules and apply rules or simple aggregation. The underlying forensic detail is gone before the scoring model sees it. Fraud Hub operates on raw, pre-binary signals. The ML engine sees every forensic data point from every domain simultaneously. This is not a feature difference. It is an architectural one.

Outsourced Signal Pipeline
Third-party module → Binary output → Rules engine → Decision. Forensic detail lost at every handoff.

Fraud Hub's Signal Pipeline
Four owned domains → Raw forensic signals → ML engine → Calibrated score. Full detail preserved.

Fraud Hub: Production-Validated

Measured in production deployment with a major financial services client. Both accuracy dimensions improved simultaneously, a result that binary architectures are unable to replicate.

70%

More Fraud Caught
Reduction in fraudulent documents accepted through verification

57%

Fewer False Rejections
Reduction in legitimate users incorrectly blocked during verification

Frequently Asked Questions

How does Fraud Hub integrate into a verification workflow?
Fraud Hub activates within the Shufti verification pipeline. For existing Shufti clients, it switches on with no additional integration. For new deployments, signal collection runs passively during the standard verification process with no changes to the user-facing flow.

What does "calibrated" mean in the context of the Risk Score?
A calibrated score means the numerical value directly represents fraud probability. A score of 0.7 corresponds to a 70% likelihood of fraud, not an arbitrary threshold or relative ranking. This calibration supports precise threshold-setting and regulatory explainability.

How does Fraud Hub differ from rule-based fraud scoring?
Rule engines combine binary flags through predefined logic. Fraud Hub's ML model ingests raw forensic signals before binary compression and detects cross-domain correlations that rule-based systems structurally cannot identify. Dynamic signal weighting adjusts automatically based on the fraud context of each session.

Does the Risk Score meet regulatory explainability requirements?
Every score decomposes into contributing signals across all four detection domains. The signal breakdown provides the audit trail and factor decomposition that compliance teams might require.

What fraud types does Fraud Hub detect that standalone modules miss?
Grey-zone fraud: coordinated attacks where each individual signal appears borderline but the combination reveals clear manipulation. This includes synthetic identities with AI-generated documents, injection attacks using replayed genuine biometrics, and mule accounts operating with clean device fingerprints on suspicious networks.



See Fraud Hub in Action

Evaluate how Fraud Hub detects the grey-zone fraud your current stack misses. Request a technical walkthrough with signal configuration for your environment.

www.shufti.com
 sales@Shufti.com

Request a Demo