
YOUR WATCH BEARS WITNESS: THE DOCTRINE OF EPISTEMIC MEDIATION AND MACHINE-GENERATED EVIDENCE IN INDIAN CRIMINAL TRIAL

Devyani Rastogi & Gaurish Sharma,
O.P. Jindal Global University, Jindal Global Law School

ABSTRACT

Evidence law has traditionally been built to test human testimony, that is, a witness's perception, memory, sincerity, and narration. But an increasing share of criminal evidence today is no longer generated by a human declarant at all. It comes from wearable devices, smart home sensors, and forensic algorithms that sense, process, and interpret physical reality before any human being ever perceives it. Under Indian law, Section 63 of the Bharatiya Sakshya Adhiniyam, 2023 treats this kind of evidence as an ordinary electronic record, verified through a dual certificate and a hash value that establishes the authenticity and integrity, but it says nothing about whether the interpretive process that produced the recorded conclusion was itself reliable.

This article argues that authenticity paradigm is conceptually inadequate to what courts are actually being asked to rely on. It also develops a Doctrine of Epistemic Mediation under which machine-generated evidence is treated as the product of a sequential chain of technological transformations, namely sensing, measurement, signal processing, software processing, algorithmic interpretation, behavioural classification, and legal characterisation, each of which introduces its own, independently assessable risk of error.

Using a hypothetical Indian prosecution modelled on the Connecticut "Fitbit murder" case, the article develops four doctrinal instruments to operationalise this theory: the Principle of Inferential Distance, the Machine Credibility Test, Algorithmic Cross-Examination, and a Digital Forensic Due Process Framework. It builds them within, rather than against, the existing scholarship of Andrea Roth, Edward Imwinkelried, and Sabine Gless.

The article also tests the theory against comparative reasoning from the United States and the European Union, against the constitutional structure of Articles 20(3) and 21 of the Indian Constitution, and against the empirical literature on wearable-sensor accuracy, before candidly examining the limits, objections, and institutional risks of the framework it proposes.

I. Introduction

On 23 December 2015, Connie Dabate was shot dead inside her Connecticut home. Her husband, Richard Dabate, told police a masked intruder had broken in. He alleged that the intruder tied him up, then assaulted him and shot Connie. There were no other eye witnesses of this incident. The Police did recover a firearm, but the forensic traces around the scene were murky at best. And Richard's account of how the intruder moved through the house didn't hold up well against the layout of the home itself. Taken on its own, none of that physical evidence was enough to prove much of anything. What actually broke the case was not a person or a document in the usual sense, but a small electronic device worn on the deceased's wrist. Connie Dabate's Fitbit fitness tracker had logged her movement through the house for roughly an hour after her husband said an intruder had already shot her. At trial, the device's data was admitted as scientifically reliable after a foundational hearing. A researcher familiar with the same model of tracker testified to its general accuracy, but candidly admitted that he did not know how the underlying algorithm converted raw sensor data into the step counts the jury was asked to rely on. Richard Dabate was convicted of murder and is currently serving a sixty five year sentence.¹

The case matters here not because it is unusual but because it is ordinary in every respect except one: the decisive witness against the accused was a machine, and the court's inquiry into that witness's trustworthiness stopped at the wrist, without ever going into the chip itself.

The expert vouched for the device's accuracy in general, but no one examined the specific algorithm that turned wrist acceleration and optical pulse readings into what the courtroom treated as “she was alive and walking.” This gap, between authenticating an output and actually understanding the process that produced it, is what this article is about.

Suppose the same case arose before an Indian sessions court. The prosecution would tender the deceased's wearable data as an electronic record under section 63 of the Bharatiya Sakshya Adhiniyam, 2023 (“BSA”),² along with a certificate under sub-section (4), one part signed by the person in operational charge of the relevant computer system and the other by an expert, and a cryptographic hash value showing that the file produced in court is identical to the file extracted from the device.

This certification architecture is inherited from Section 65B of the Indian Evidence Act, 1872, and is now codified in substantially similar form.³ It answers a narrow question: is this the same record, unaltered, that the device actually generated? That question is worth asking, but it isn't the same as asking whether the process behind the record was reliable, whether the way a device converted

¹ State v Dabate, 351 Conn 428 (2025).

² Bharatiya Sakshya Adhiniyam 2023, s 63.

³ Indian Evidence Act 1872, s 65B.

a physical event into a behavioural label, such as "walking," "running," or "under stress," was itself trustworthy. Section 63 was never designed to answer that second question.

This article's central claim is that the second question is anything but a minor footnote to the law of electronic evidence. It is, this article argues, the central epistemological problem that machine-generated proof poses for criminal adjudication. Existing doctrine, in India and elsewhere, simply has no settled vocabulary for asking it. Wearable devices serve as the main case study here for two reasons. They are already before Indian courts, in matrimonial disputes, insurance claims, and increasingly in criminal cases. And their sensing-to-inference architecture is well documented in the biomedical engineering literature, which means the theoretical claims made below can be tested against real data instead of guesswork. But the argument itself isn't really about wearables. It reaches the wider category of machine-generated evidence: smart-home systems, autonomous-vehicle logs, medical implants, drone and CCTV analytics, biometric and predictive forensic software. Wearables simply offer the most tractable, best-studied example of that category.

This article proceeds in fourteen parts. Part II locates the conceptual gap in existing doctrine. Part III examines, rather than assumes, the increasingly common label "machine witness." Part IV asks what a legal fact actually is when the fact in question is an algorithm's inference. Part V develops the philosophical architecture of the argument, a seven-stage Knowledge Production Chain grounded in the philosophy of technology and evidence epistemology, and tests it against the empirical literature on wearable-sensor error. Part VI sets out the Doctrine of Epistemic Mediation and its four instruments. Part VII separates out seven concepts of judicial trust that existing doctrine tends to conflate. Part VIII proposes a taxonomy of machine-generated proof. Part IX addresses subsidiary doctrinal questions of hearsay, trade secrecy, burden allocation, and differential standards of proof. Part X undertakes a comparative analysis of judicial reasoning in the United States, the European Union, and India. Part XI tests the theory against the existing literature and states plainly what is, and is not, original about it. Part XII examines the theory's limits and likely objections. Part XIII proposes an implementation pathway for Indian courts, and Part XIV concludes.

II. The Problem: From Testimony to Technological Mediation

A. What the Traditional Approach Tests

The architecture of the law of evidence, in India as elsewhere, was built around one recurring problem: an aggrieved person reports something to the court, and the court has to decide whether there exists enough evidence to support their claims. The four traditional testimonial infirmities, defects of perception, memory, narration, and sincerity, are all properties of a human declarant. Cross-examination exists to test exactly these properties: did the witness actually see what she claims to have seen, has her memory of the event decayed or been contaminated, is her account of

the event accurate, and is she telling the truth as she believes it or lying. Hearsay is excluded, subject to numerous exceptions, mainly because an out-of-court declarant who is not before the court cannot be tested on any of these four dimensions.

None of this machinery is well suited to a wearable device like a smart watch, a smart home hub, or a forensic classification algorithm. A device does not perceive in the human sense, it transduces a physical signal. It does not remember, it stores. It does not narrate, it simply outputs a value or a label according to a fixed or trained procedure. And it cannot lie, because it has no beliefs to misrepresent in the first place. Yet a wearable's classification of a wearer's cardiovascular state as "stressed," or a smart lock's log of a door opening, still functions in the courtroom exactly as human testimony once did: as an assertion offered for the truth of what it asserts, which the fact-finder is invited to rely on.

B. The Authenticity Paradigm and Its Limits

Section 63 of the BSA, like its predecessor section 65B of the Indian Evidence Act, addresses this problem by simply re-describing the machine's output as a document. A "computer output" is deemed to be a document, and is therefore admissible without further proof of what it states, if it satisfies four conditions: regular use of the computer, the ordinary course of the activities carried on, the proper operation of the computer during the relevant period, and the accuracy of the information fed into it. The Supreme Court of India has treated this certification requirement as mandatory rather than merely directory.

In the case of *Anvar P.V. v. P.K. Basheer*, the Court held that Section 65B lays down a complete code for admitting electronic records⁴. Without the statutory certificate, secondary evidence of an electronic record simply could not be led. For a while, courts chipped away at that rule. One line of cases said the certificate requirement could be relaxed if getting the certificate was genuinely impossible. That relaxation didn't survive. A larger bench in the case of *Arjun Panditrao Khotkar*

v. Kailash Kushanrao Gorantyal stepped in, reaffirmed the holding given in the *Anvar* case, and made clear that the certificate is a condition precedent to admissibility with only a narrow exception carved out.⁵

This body of law is rigorous and it does an important job: it guards against tampering, misattribution, and the silent substitution of one record for another. But every one of its tests is directed at the question of authenticity, meaning is this the record the device actually produced, and at the question of proper operation at a general, systemic level, meaning was the computer functioning normally during the relevant period. Neither the statutory text nor the case law asks

⁴ *Anvar PV v PK Basheer* (2014) 10 SCC 473.

⁵ *Arjun Panditrao Khotkar v Kailash Kushanrao Gorantyal* (2020) 7 SCC 1.

whether the specific interpretive process by which the device converted raw sensor readings into the label now presented as evidence was itself a sound process.

A hash value proves that a file has not been altered since it left the device. However, it cannot prove that the algorithm which produced the file's contents was accurate, well calibrated, validated against a representative population, or free from the kind of context-dependent error that the biomedical engineering literature, discussed in Part V, shows to be common in exactly this class of device.

The conceptual deficiency, in other words, is a category mistake. Existing doctrine treats the reliability of a machine-generated conclusion as if it were reducible to the integrity of a file. Integrity and reliability are related but different properties, and a legal test built for the first cannot, on its own, discharge the second. The rest of this article develops a doctrine addressed specifically

to the second problem, not to displace the authentication framework of section 63, which remains necessary, but to supply the inferential-reliability inquiry that framework does not, and structurally cannot, perform.

III. Is "Machine Witness" the Right Concept? A Critical Examination

The phrase "machine witness" has begun to appear in judicial and academic writing as shorthand for evidence that originates from a device rather than a person. It is an appealing label, largely because it renders something unfamiliar in terms we already understand but convenience is not the same as accuracy, and a label adopted for convenience risks smuggling in assumptions that do not hold up under scrutiny. Before this article makes use of the term, four candidate framings need to be tested against what a device actually does.

The first option is the witness analogy itself. *Andrea Roth*, its strongest defender, points out that many machines now generate statements no programmer or operator could have independently verified at the time, some of them with a testimonial character to them.⁶ Courts, faced with this tend to improvise. One might treat the output as hearsay. Another might treat it as real evidence. A third might treat it as the raw data behind a human expert's opinion.

That doctrinal improvisation is real, and Roth rightly points out that some machine outputs are credibility dependent in a way that resembles testimony. But the analogy breaks down at the point that matters most here: a witness's trustworthiness depends on sincerity, on the absence of any motive to lie. A device has no sincerity to assess, no motive to interrogate. The witness analogy is useful for one reason: it reminds a court that trustworthiness needs to be assessed before an assertion can be relied on. But it obscures something just as important: what that assessment actually consists of is not the same thing at all.

⁶ Andrea Roth, 'Machine Testimony' (2017) 126 Yale LJ 1972.

The second option is the scientific instrument. A wrist-worn optical sensor is, at its first stage, genuinely an instrument. It converts a physical quantity, specifically light absorption in blood vessels, into an electrical signal. A thermometer does something similar, converting thermal expansion into a needle position. The instrument analogy holds for that first stage. It breaks down for everything after it. A thermometer measures. It doesn't interpret.

However, a wearable's downstream software does interpret. It classifies a pattern of processed signal as "running" or as "stressed." A simple measuring instrument has no equivalent step. And that step introduces a different category of error entirely, one discussed in Part V.

The third option is to call it a document, or an electronic record. This is actually the label Indian law has gone with. It works well for capturing things like storage and reproducibility, which is why Section 63's certification rules hold up reasonably well for that purpose. The problem is that this framing treats the machine's output like a static piece of text, when it isn't one. That's the same mistake flagged back in Part II. A document just holds a proposition someone already made. A classification algorithm is different. It's the one generating the proposition in the first place.

The fourth option is to treat the machine as something new entirely, an "autonomous observer," its own category of evidence. Of the four options, this one is the most honest about what's actually happening. But it comes with a real risk: unless it's built on top of existing rules for admissibility, disclosure, and appeals, it could end up too vague to actually use in court.

Here's where this article lands. "Machine witness" is a useful phrase, but only as a shorthand for one idea: that trusting a machine's output requires the same kind of upfront trust judgment that trusting a human witness does. It shouldn't be taken to mean the device is literally testifying, or that it has anything like a witness's capacities. For that reason, the term used throughout the rest of this article is machine source instead. "Witness" is set aside for a narrower purpose, the credibility-testing framework laid out in Part VI. None of this conflicts with the American case law holding that a machine can't be a hearsay declarant, since hearsay requires a person making an assertion. If anything, this article's argument depends on that same idea.

In the case of *United States v. Lizarraga-Tirado*, the Ninth Circuit held that computer generated map coordinates were not hearsay precisely because no person made the relevant assertion⁷. That reasoning has been applied to other categories of machine-generated data as well. This is correct on its own terms, a machine is not a declarant, but it is incomplete, because it removes machine output from the one doctrinal mechanism, hearsay, and the cross-examination it presupposes, designed to test a source's credibility, without putting anything in its place. The "no declarant" holding closes a door. It does not open the one this article argues is needed.

⁷ *United States v. Lizarraga-Tirado*, 789 F3d 1107 (9th Cir 2015).

IV. Jurisprudential Foundation

The central argument of this article is built on a distinction that evidence scholars tend to acknowledge only in passing, rather than examine directly: the distinction between observing something and inferring it. In most contexts, this difference doesn't carry much practical weight.

Here, it does because the "observer" in question is a machine and machines do not merely record what happens, they also draw conclusions, often within the same step as the observation itself, and that shift from seeing to inferring tends to pass unnoticed.

A. What Is a Legal Fact?

For adjudicative purposes, a fact is usually understood as something direct evidence could prove, something a person could have seen or heard and then reported. A raw sensor voltage, a GPS coordinate, a timestamp: these are all close to that idea of observation. They're roughly what a witness would describe seeing or hearing, just picked up by a machine instead of a person.

B. Does Machine Generated Evidence Discover Facts, or Construct Them?

A smartwatch's conclusion that its wearer was "running," or was in a state of "high stress," is not an observation in this sense at all. It is a classification, the product of a trained or hand-coded model applied to processed sensor data, choosing among a fixed vocabulary of behavioural labels the one judged statistically most probable.

No instrument actually observes "stress." What is measured is heart-rate variability, electrodermal activity, or a similar physiological correlate.

Similarly, "Stress" is a label that the model assigns to a pattern in that measurement. It is built into the model, carrying contestable assumptions about what stress actually is, and about how it's supposed to show up in the body. In *Susan Haack's* terms, the output is really a knowledge claim, a probabilistic, model relative proposition, not the kind of fact the law of evidence has traditionally assumed it to be.

Haack makes a broader point in her critique of American reliability jurisprudence: courts tend to confuse being "scientific" with being reliable⁸. The label itself gets treated as proof, when it isn't. Something similar happens with machines. Courts often treat a machine's output as self-authenticating just because the device looks technically sophisticated. But a number isn't more trustworthy because a microprocessor produced it instead of a person. It's only as trustworthy as the process behind it that hasn't changed, and it isn't going to.

⁸ Susan Haack, *Evidence Matters: 'Science, Proof, and Truth in the Law'* (Cambridge University Press 2014).

C. The Doctrinal Consequence: Distinguishing Observation from Interpretation

The consequence for evidence law is that it should distinguish, far more sharply than it currently does, between raw observations and interpreted conclusions, and should treat the latter as carrying an accumulated inferential risk that the former does not.

This is not just a semantic point. It determines how far a court's inquiry has to travel before it can responsibly rely on a machine-generated proposition: the further an output sits from direct transduction of a physical quantity, the more interpretive judgment has been built into it, and the more its presentation as a settled "fact" is a legal fiction that needs justification rather than a starting assumption.

Part V develops the architecture that operationalises this distinction, and Part VI builds the doctrine on top of it.

V. The Philosophy of Proof: The Knowledge Production Chain

A. Technological Mediation is Never Neutral

This article's theory begins with the philosopher *Don Ihde*, who studied how instruments shape the way human beings perceive the world. His central claim is that no instrument presents reality directly. Every instrument reveals certain features while concealing others, and the two effects cannot be separated. For example, a pair of eyeglasses magnifies vision, yet narrows what can be seen at the periphery. Any instrument that brings one feature into focus does so by pushing another out of focus. Ihde further observes that instrument readings are not self-explanatory. They require interpretation, and only a reader with the relevant training can interpret them correctly.

This has real consequences for courts. When a court treats a wearable's printout as a direct, unmediated view of what the wearer's body was doing, it makes a mistake. The device reveals and conceals at the same time, and reading it correctly requires expertise the court doesn't have unless someone supplies it. The printout isn't a record of what happened. It's a record of what a chain of sensors, filters, and models concluded had happened, shaped by design choices that were never placed before the court.

B. Reliability is Not Reducible to a Single Judgment

The second foundation comes from the philosopher *Susan Haack*. She argues that trustworthy knowledge doesn't rest on a single foundational fact. It's built more like a crossword puzzle: many

9 Don Ihde, *Technology and the Lifeworld: 'From Garden to Earth'* (Indiana University Press 1990).

small pieces that check and support one another, rather than a tower balanced on one base.

¹⁰That image suits this article's purpose, since what's being judged here isn't a single output but an entire chain of steps. No individual step in a machine's process needs to be perfect on its own. What matters is how the steps work together, checking and reinforcing each other. A weakness in one step can quietly propagate through the rest of the chain without surfacing at the end, which is usually the only part a court actually sees.

C. The Seven Stage Knowledge Production Chain

These two ideas, from Ihde and Haack, set up the rest of this Part.

Ihde's argument: every instrument shows you something by hiding something else, so a device's output is never a neutral window onto reality. Haack's argument: trustworthy knowledge comes from a chain of steps checking each other, not from one perfect step. Together, they point to the same conclusion.

If a device's output is shaped at every stage, and if trustworthiness depends on the whole chain rather than any single link, then machine-generated evidence should be evaluated stage by stage, not treated as one single, all or nothing fact.

This article proposes a seven stage architecture for doing exactly that, using a wrist worn wearable (smart watch) as the running example throughout.

1. Sensing / Transduction

A physical quantity, say, light absorption in blood vessels beneath the skin, or the acceleration of a wrist through space, gets converted into an electrical signal by a physical sensor. Errors creep in early here: sensor placement, contact quality, ambient light interference. Skin tone matters too because the green light photoplethysmographic sensors responds differently depending on melanin content, which raises both an accuracy problem and a fairness problem, taken up again in Part XII.¹¹

2. Measurement / Quantisation

The device cannot record a continuous signal forever, so it takes snapshots of it at regular intervals and converts each one into a digital number. That process, called quantisation, introduces its own errors.

If the snapshots are taken too infrequently, or if the digital conversion isn't precise enough, or if the signal exceeds what the sensor can actually capture, the resulting numbers can already be off before any further processing even begins.

¹⁰ Susan Haack, *Evidence and Inquiry: Towards Reconstruction in Epistemology* (Blackwell 1993).

¹¹ Brinnae Bent, Benjamin A Goldstein, Warren A Kibbe and Jessilyn P Dunn, 'Investigating Sources of Inaccuracy in Wearable Optical Heart Rate Sensors' (2020) 3 NPJ Digital Medicine 18.

3. Signal Processing / Filtering

Once digitised, the signal is still messy. It gets smoothed and cleaned up to strip out motion artefacts and isolate the physiological signal underneath. This step is harder than it sounds, because the software has to assume, in advance, what counts as "noise" and what counts as "signal," and that assumption isn't always right. A filter tuned for someone sitting still can misfire once that person starts moving. It may suppress a signal that's actually real, or invent one that isn't there.

This isn't hypothetical. Wrist worn optical sensors have a well known failure called "signal crossover."¹² Limb motion, arm swing during walking, for example, produces its own strong, repetitive rhythm. The sensor can mistake that rhythm for the heartbeat it's meant to be measuring. When that happens, the heartrate reading can be wrong at exactly the moment of activity most likely to matter in court.

4. Software Processing

Sensor data does not reach the classifier untouched. The device's firmware processes it first by smoothing the stream, removing readings that fall outside expected ranges, converting the values into standard units. These are ordinary operations, but ordinary does not mean reliable. A defect in the processing code can alter the output, and so can a firmware update, which may install automatically and without any record of what it changed. Should the device later become evidence, the situation is stark: neither party may be aware that the software was ever modified, and neither is positioned to determine whether that modification affected the result.

5. Algorithmic / Statistical Interpretation

After the signal is processed, an algorithm, often based on machine learning converts it into a measurable output, such as heart rate in beats per minute. Its accuracy depends on how the model was designed and whether it was trained or validated on a population similar to the user. Performance also varies with activity.

One study found a mean absolute percentage error of only 1.2% during running, but 16–17% during badminton and football, where rapid arm movements interfere with sensor readings¹³. Another study reported that wrist worn devices were least accurate during sudden movement or rapid changes in heart rate, compared to periods of steady activity.¹⁴

¹² Bent, Goldstein, Kibbe and Dunn (n 11).

¹³ Maxime Ceugniet, Hervé Devanne and Eric Hermand, 'Reliability and Accuracy of the Fitbit Charge 4 Photoplethysmography Heart Rate Sensor in Ecological Conditions: Validation Study' (2025) 13 JMIR mHealth and uHealth e54871.

¹⁴ Catharina Nina Van Oost and others, 'Accuracy of Heart Rate Measurement Under Transient States: A Validation Study of Wearables for Real-Life Monitoring' (2025) 25 Sensors 6319.

6. Behavioural Classification

The processed data is then assigned a behavioural label like "running," "asleep," "highly stressed" by applying predefined thresholds or a classification algorithm. The reliability of that label turns on several things: where the thresholds were set, how finely the categories distinguish one state from another, and how the underlying model was developed. A classifier well suited to everyday fitness tracking may nonetheless produce too many false positives to be safe in a criminal trial.

7. Legal Characterisation

Presented in court, this output is often treated as evidence of a legal fact: that the deceased was alive and moving at a given moment, for instance. Yet the uncertainty built up across the earlier stages has, by now, dropped out of sight. What reaches the court looks like a settled factual conclusion, when in truth it is the outcome of several layered probabilistic judgments.

D. What the Chain Shows

This evidentiary chain exposes two weaknesses in how machine generated evidence is currently assessed.

The first concerns reliability, which does not hold constant across the process. A wearable may record physiological data accurately while the software that processes and classifies that data introduces uncertainty of its own. To ask only whether the device is "accurate" is to miss the stages where the errors actually arise.

The second weakness runs deeper. As the empirical studies show, wearables perform least accurately during rapid movement, physical exertion, and sudden physiological change at precisely the moments most likely to matter when a criminal event is reconstructed, be it a struggle, a flight, or a collapse.

Section 63 of the BSA nevertheless directs its attention to authenticating the final output, not to the reliability of the stages that produced it. Part VI accordingly proposes a stage specific approach that the courts must assess the reliability of each step in the chain rather than treat the final output as trustworthy in itself.

VI. The Doctrine of Epistemic Mediation

A. Statement of the Core Doctrine

This article proposes the following principle for evaluating machine generated evidence in criminal proceedings:

Doctrine of Epistemic Mediation: Machine-generated evidence does not directly represent reality. It is the end product of a sequence of technological processes through which physical events are converted into digital information and, ultimately, into legal proof. Courts should therefore assess the reliability of each stage of that process, rather than treating the final output as reliable simply because it appears authentic.

The central idea is straightforward. Reliability cannot be determined by examining only the final output. Every stage involved in producing that output has the potential to introduce error, and each should therefore form part of the court's assessment.

This is not an entirely new insight. As Part XI explains, both Roth's account of machine credibility and Imwinkelried's work on foundational reliability recognise that machine evidence requires more than simple authentication.¹⁵

This article builds on those ideas in a different way. It develops a seven stage framework based on engineering research on error propagation from data collection to algorithmic inference. Unlike Roth's broader typology, or the lifecycle models used in digital forensics, the framework focuses on how the meaning of machine-generated evidence is produced and where reliability may be lost before the evidence reaches the courtroom. The remainder of this Part explains four principles through which the doctrine can be applied in practice.

B. The Principle of Inferential Distance

The first principle translates the Knowledge Production Chain into a guide for evaluating evidentiary weight.

Principle of Inferential Distance: The further a machine-generated conclusion is removed from the original sensor reading, the less evidentiary weight it should receive unless it is supported by independent evidence.

Not every machine-generated output deserves the same degree of confidence. A raw timestamp or an unprocessed GPS coordinate has undergone relatively little interpretation and may therefore carry considerable evidentiary weight once authenticated. By contrast, a conclusion such as "highly stressed" represents the end of a much longer chain of processing, modelling, and

¹⁵ Roth (n 6); Edward J Imwinkelried, 'The Admissibility of Scientific Evidence: Exploring the Significance of the Distinction Between Foundational Validity and Validity as Applied' (2020) 70 Syracuse L Rev 817.

classification. Without independent support such as medical records, witness testimony, or other forensic evidence, its evidentiary value should be treated more cautiously.

This is not intended to create a new rule of admissibility. Rather, it offers a structured way for courts to assess the weight of machine-generated evidence. Just as the law distinguishes between direct and indirect testimony, it should also recognise that different machine-generated conclusions involve different degrees of interpretation before they reach the courtroom.

C. The Machine Credibility Test

The doctrine also requires a structured method for assessing whether a particular chain of technological mediation is reliable enough to support the inference for which it is offered.

Machine Credibility Test: In assessing the reliability of machine-generated evidence, courts should consider, so far as practicable: (i) the accuracy of the relevant sensor and its documented error rates for the activity in question; (ii) the device's calibration history; (iii) the integrity of its firmware and software, including version and update history; (iv) its cybersecurity safeguards and vulnerability to tampering; (v) the transparency of the algorithm and the extent to which its classification process can be explained; (vi) the existence and quality of independent, peer-reviewed validation studies; and (vii) the integrity of any cloud synchronisation or data transmission pathway.

The factors are drawn from established scholarship on machine-generated evidence but are reorganised around the seven-stage framework developed in this article. They also reflect the distinction made by the President's Council of Advisors on Science and Technology between foundational validity, which asks whether a method is capable of producing reliable results in principle, and validity as applied, which examines whether that method was used reliably in the particular case.¹⁶ Sensor accuracy, calibration, and independent validation primarily address foundational validity. Firmware integrity, cybersecurity, and data transmission are concerned with validity as applied, while algorithmic transparency is relevant to both.

D. Algorithmic Cross-Examination

Reliability cannot be assessed unless the defence has a meaningful opportunity to examine how the machine reached its conclusion. Unlike a human witness, however, a machine cannot be questioned in court. The scrutiny must instead focus on the technological process that produced the evidence.

¹⁶ President's Council of Advisors on Science and Technology, *Forensic Science in Criminal Courts: Ensuring Scientific Validity of Feature-Comparison Methods* (Executive Office of the President, September 2016).

Algorithmic Cross-Examination: The procedural mechanism through which a party challenges machine-generated evidence by obtaining pre-trial disclosure of the relevant device or system's design documentation, validation studies, and information necessary for independent testing, including source code or model parameters where appropriate, together with the opportunity to call expert evidence on the stages of the mediation chain placed in issue.

This proposal builds on *Christian Chessman's* argument that criminal defendants should be able to examine the source code of software used to generate evidence against them¹⁷. It also reflects *Rebecca Wexler's* work showing that trade-secret claims should not operate as an evidentiary privilege that prevents meaningful scrutiny¹⁸. The expression "algorithmic cross-examination" is therefore shorthand for a process of disclosure, independent testing, and expert evaluation. It does not suggest that an algorithm can be questioned in the same way as a witness.

E. The Digital Forensic Due Process Framework

The doctrine also requires procedural safeguards at every stage of the evidentiary process, from the seizure of the device to the court's final assessment of its evidentiary value. These safeguards include authentication under section 63, forensic preservation techniques that avoid altering the device, disclosure of software and model versions, independent validation through the Machine Credibility Test, and an unbroken chain of custody. Courts should also provide reasoned findings on any disputed stage of the mediation chain instead of merely noting that a section 63 certificate has been produced. This framework is distinct from the lifecycle models used in digital forensics. While those models are designed to preserve the integrity and custody of digital evidence, they do not examine whether the processes that transformed raw sensor data into the final evidentiary output were themselves reliable.

¹⁷ Christian Chessman, 'A "Source" of Error: Computer Code, Criminal Defendants, and the Constitution' (2017) 105 Cal L Rev 179.

¹⁸ Rebecca Wexler, 'Life, Liberty, and Trade Secrets: Intellectual Property in the Criminal Justice System' (2018) 70 Stan L Rev 1343.

VII. A Taxonomy of Machine-Generated Proof

Treating all electronic records as a single category overlooks the different risks associated with different forms of machine-generated evidence. The reliability concerns raised by a pacemaker, for example, are not the same as those raised by a smartwatch or a facial-recognition system. To reflect these differences, this article groups machine-generated evidence into five broad categories based on the nature of the technology and the kinds of errors to which it is most susceptible. The categories are not mutually exclusive, and a single device may fall within more than one of them.

1. Passive Systems

Medical devices such as pacemakers, continuous glucose monitors, and implantable cardiac defibrillators generally have a strong evidentiary foundation because they are subject to extensive clinical testing and regulatory oversight. Their principal legal concern is therefore not measurement accuracy but the use of sensitive medical data in criminal proceedings. Questions of privacy, consent, and the privilege against self-incrimination become particularly significant when information collected for treatment is later relied upon as evidence. These issues are examined further in Part IX.

2. Environmental Systems

Smart home hubs, voice assistants, and networked security cameras record information about the surrounding environment rather than about a specific individual. That distinction creates different evidentiary concerns. A smart lock can establish that a door was opened, but not necessarily who opened it. Cloud storage and continuous synchronisation also extend the chain of custody, increasing the possibility that data may be altered, lost, or retained selectively.

3. Mobile Systems

Wearables and smartphones form the central focus of this article. Because these devices are usually associated with a single user, they can provide highly individualised evidence. At the same time, as discussed in Part V, they are especially vulnerable to errors during movement, physical exertion, and other transient physiological states. Their evidentiary value therefore depends heavily on the context in which the data was generated.

4. Predictive Systems

Risk assessment tools, facial-recognition software, and probabilistic forensic matching systems occupy the furthest end of the inferential-distance spectrum. Their outputs are shaped by multiple layers of processing and statistical modelling, making them both highly interpretive and difficult to scrutinise. As a result, concerns about transparency and meaningful challenge are particularly acute. These issues are considered in greater detail in Part VIII.

VII. Doctrinal Implications

A. Machine-Generated Evidence and the Hearsay Rule

Under the American evidence law, machine-generated output is generally treated as non-hearsay because the hearsay rule applies only to statements made by a human declarant. Although this position is doctrinally sound, it leaves an important gap. By placing machine-generated evidence outside the hearsay rule, the law also excludes it from the framework traditionally used to examine the credibility of the source. No comparable mechanism has developed to assess the reliability of the machine or the process that produced its output.

This article does not argue that machine-generated evidence should be brought within the hearsay rule. Instead, it argues that courts need a separate method for testing its reliability. The Machine Credibility Test and Algorithmic Cross-Examination are proposed to perform that role by directing judicial scrutiny towards the process through which the evidence was generated, rather than treating the final output as sufficient in itself.

B. Trade Secrets and the Right to a Fair Trial

One of the main barriers to Algorithmic Cross-Examination is the use of trade-secret claims to restrict access to source code and model parameters. This issue has received considerable attention in the United States, particularly in litigation involving probabilistic genotyping software. *Rebecca Wexler* argues that intellectual property rights should not prevent a defendant from examining the process that produced the evidence relied upon by the prosecution.¹⁹ The same principle should apply in India. The guarantee of a fair trial under Article 21 supports disclosure of material necessary to test the reliability of the evidence, while allowing courts to protect genuinely confidential information through appropriate protective orders.²⁰

C. Understanding Algorithmic Cross-Examination

The expression "algorithmic cross-examination" should not be understood literally. An algorithm cannot be questioned in the same way as a witness. Instead, the term describes a procedural mechanism through which an opposing party can test the reliability of machine-generated evidence. This includes access to relevant technical material, independent testing where necessary, and expert evidence directed at the stages of the mediation chain that are genuinely in dispute. The objective is not to question the machine itself but to examine the process that produced its output.

¹⁹ Wexler (n 18).

²⁰ Constitution of India 1950, art 21.

D. The Burden of Establishing Reliability

Where the reliability of machine-generated evidence is challenged, the burden should fall on the party seeking to rely upon it, which in criminal proceedings will ordinarily be the prosecution. This approach is consistent with the prosecution's obligation to establish guilt beyond reasonable doubt. It also follows naturally from section 63, which already requires the party tendering the electronic record to establish its authenticity. Extending that obligation to include the reliability of the process that generated the evidence represents a modest but important development.

E. Different Categories, Different Standards

The taxonomy proposed in this article recognises that not all machine-generated evidence raises the same concerns. The level of scrutiny should therefore reflect both the category of evidence and its inferential distance. Passive systems, environmental systems, and low-inference outputs from mobile devices, such as raw timestamps or unprocessed location data, will often be capable of standing on authentication alone. Predictive systems, hybrid systems, and high-inference outputs from mobile devices, including behavioural and physiological classifications, require greater caution because they are further removed from the underlying sensor data. As a general rule, they should be supported by independent corroborative evidence before they are relied upon to establish guilt. That expectation may be displaced only where the Machine Credibility Test demonstrates that the evidence is sufficiently reliable.

IX. Comparative Analysis: Judicial Approaches Across Jurisdictions

Different legal systems have responded to machine-generated evidence in different ways. Their approaches reflect different priorities. Some focus on scientific reliability, others on regulatory oversight, while Indian law has largely concentrated on the authenticity of electronic records. Examining these approaches together helps identify both the strengths and the limits of the current Indian framework.

A. The United States of America

In the United States, questions about machine-generated evidence are addressed through two distinct bodies of law. Scientific reliability is assessed under *Daubert v. Merrell Dow Pharmaceuticals*.²¹ The constitutional right to confront adverse evidence, on the other hand, is shaped by *Crawford v. Washington*²² and later decisions such as *Melendez-Diaz v. Massachusetts*.²³

21 *Daubert v Merrell Dow Pharmaceuticals Inc*, 509 US 579 (1993).

22 *Crawford v Washington*, 541 US 36 (2004).

23 *Melendez-Diaz v Massachusetts*, 557 US 305 (2009).

The distinction becomes important when evidence is generated by opaque or “black-box” software. A defendant may challenge the scientific reliability of the technology under the Daubert reasoning. Yet the Confrontation Clause provides no equivalent opportunity to question the process that produced the result because an algorithm is not a witness. American law therefore offers a well-developed framework for evaluating scientific reliability but provides far less guidance on how the underlying technological process should be scrutinised. The model of Algorithmic Cross-Examination proposed in this article seeks to address that gap through disclosure, independent testing, and expert evaluation.

B. European Union

The European Union has taken a different approach. Rather than leaving questions of reliability to be resolved during litigation, the Artificial Intelligence Act regulates certain AI systems before they are deployed.²⁴ AI systems used by law enforcement to assess the reliability of evidence are classified as high-risk and must satisfy detailed requirements relating to risk management, documentation, transparency, and human oversight.

This shifts much of the inquiry away from the courtroom because by requiring extensive documentation before deployment, the AI Act promotes greater consistency across systems. Even so, it says comparatively little about the weight a judge should attach to a particular output in an individual case. The framework proposed in this article complements that model. Documentation prepared to satisfy the AI Act can also assist courts when applying the Machine Credibility Test.

C. India

Indian law has developed a rigorous framework for establishing the authenticity of electronic records. Section 63 of the Bharatiya Sakshya Adhiniyam, together with its predecessor section 65B of the Indian Evidence Act, has been interpreted in the cases of *Anvar* and *Arjun Panditrao* as making the statutory certificate a mandatory requirement for admissibility. The Bharatiya Nagarik Suraksha Sanhita, 2023 has also expanded investigative powers relating to electronic records.²⁵

Authenticity, however, is only part of the inquiry. A section 63 certificate may establish that an electronic record accurately reproduces the data stored on a device. It does not establish that the conclusions drawn from that data are correct. A wearable device might classify a person’s activity as “walking” when the individual was in fact being dragged. The record may be authentic even though the inference is wrong. No reported Indian appellate decision has yet examined this issue

24 Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act) [2024] OJ L1689

25 Bharatiya Nagarik Suraksha Sanhita 2023.

in the context of behavioural inferences generated from wearable devices. The opening illustration in this article is intended to anticipate that question before it emerges in litigation.

The constitutional framework reflects a similar distinction. Article 20(3) of the Indian Constitution protects the accused from forced self incrimination, though the scope of that protection has never been self-evident from the text, it is judicial interpretation that has done most of the work²⁶. *Selvi v State of Karnataka* is the clearest example²⁷. There, the Supreme Court held that the guarantee could not be confined to conventional confessions, and extended it to investigative techniques such as narco-analysis, polygraph testing, and Brain Electrical Activation Profile (BEAP) testing, on the reasoning that each of these techniques extracted communicative information without the subject's voluntary participation. The Court stopped short of collapsing an older distinction, though. *State of Bombay v Kathi Kalu Oghad* had already separated testimonial compulsion from the compelled production of physical or identifying material, and *Selvi* left that line intact²⁸. So, Article 20(3), as it stands, protects compelled communication specifically, not every form of evidence a court might obtain from an accused person.

Privacy presents another, separate challenge. Since the *K.S. Puttaswamy* case, Article 21 has been understood to protect personal privacy, subject to the requirement of proportionality.²⁹ Wearable devices often record continuous streams of information extending well beyond the period relevant to a criminal investigation. The question is therefore not only whether the data is reliable, but also whether the scope of forensic extraction is proportionate. The Digital Forensic Due Process Framework proposed in this article is intended to operate within those constitutional limits.

X. Limits, Objections, and Unintended Consequences

No doctrinal framework is without limitations. The proposal advanced in this article is no exception. Several objections require careful consideration. This is because they expose the practical limits of the framework and also identify issues that any future implementation would need to address.

A. The Risk of Over-Inclusiveness

If every item of machine-generated evidence were subjected to stage-by-stage scrutiny, the framework could become unnecessarily burdensome. Routine evidence, such as a server timestamp, would rarely justify that level of examination, and applying the full framework in every case would place additional pressure on already congested trial courts. The Principle of Inferential Distance is intended to avoid that result. Outputs involving only limited interpretation should ordinarily require nothing more than authentication under section 63. The more demanding

²⁶ Constitution of India 1950, art 20(3).

²⁷ *Selvi v State of Karnataka* (2010) 7 SCC 263.

²⁸ *State of Bombay v Kathi Kalu Oghad* AIR 1961 SC 1808.

²⁹ *Justice KS Puttaswamy v Union of India* (2017) 10 SCC 1.

Machine Credibility Test is reserved for evidence that has passed through multiple interpretive stages.

B. Institutional Capacity

The framework also assumes a level of technical expertise that many trial courts may not possess. Questions about firmware integrity, calibration records, or model validation often require specialised knowledge. There is also a risk that the proposed framework could become another procedural formality, much like the section 65B certificate sometimes did in practice. Existing institutions offer a partial answer. The Examiner of Electronic Evidence, recognised under section 79A of the Information Technology Act, 2000, can provide technical assistance where necessary.³⁰ Courts should also explain their conclusions on each disputed stage of the mediation chain instead of treating the exercise as a routine procedural requirement.

C. The Importance of Corroboration

A further concern is that a preference for corroboration may lead courts to discount reliable evidence simply because no independent source is available. That concern is legitimate, but it does not require abandoning the framework. Corroboration is proposed as a general expectation rather than an inflexible rule. Where the Machine Credibility Test demonstrates a high level of reliability, a court should remain free to rely on a single source of machine-generated evidence.

D. Unequal Access to Information

The proposed framework cannot eliminate every practical obstacle. Technology companies may continue to rely on trade-secret claims to restrict access to source code and other technical information. Unless the law recognises that a defendant's right to a fair trial takes priority in appropriate cases, Algorithmic Cross-Examination may exist in theory while remaining difficult to exercise in practice. Any implementation of the framework should therefore address disclosure directly instead of leaving the issue to be resolved case by case.

E. Automation Bias

Even a well-designed legal framework may not prevent judges or juries from placing excessive trust in machine-generated evidence, research on automation bias shows that decision-makers often give greater weight to automated outputs than to conflicting evidence.³¹ This is exactly where procedural safeguards become important. Mandatory disclosure is one part of it. So is placing the burden on the party relying on the evidence to establish its reliability. And judicial findings that

³⁰ Information Technology Act 2000, s 79A.

³¹ Kate Goddard, Abdul Roudsari and Jeremy C Wyatt, 'Automation Bias: A Systematic Review of Frequency, Effect Mediators, and Mitigators' (2012) 19 Journal of the American Medical Informatics Association 121

are actually reasoned, rather than assumed, push toward closer scrutiny instead of unquestioning acceptance.

F. Alternative Approaches

The framework proposed in this article is not the only possible solution. One alternative would be to incorporate Daubert style reliability review through a purposive interpretation of the Bharatiya Sakshya Adhiniyam. Another would be to adopt Roth's model of testimonial safeguards. A third would require the use of interpretable or "glass-box" systems, following the approach suggested by Garrett and Rudin³². These approaches need not compete with the Doctrine of Epistemic Mediation. Instead, they address different aspects of the same problem. The doctrine proposed here provides a broader conceptual framework by distinguishing between observation and inference and by recognising that different machine-generated outputs involve different levels of interpretation before they reach the courtroom.

XI. The Proposed Framework

The principles developed in this article can be incorporated into existing evidentiary practice without waiting for legislative reform. Much of the proposed framework can be implemented through a purposive interpretation of the Bharatiya Sakshya Adhiniyam, particularly its provisions governing electronic records and expert evidence.

The starting point remains unchanged. Courts should continue to authenticate electronic records under section 63 by requiring the statutory certificate, the applicable hash values, and compliance with the existing statutory requirements. The framework proposed in this article supplements that process; it does not replace it.

The next question is the nature of the evidence before the court. Using the taxonomy developed in Part VII, judges should identify the category into which the machine-generated evidence falls and determine its inferential distance. Evidence that involves little interpretation, such as raw timestamps or unprocessed location data, will usually require no further inquiry unless its reliability is specifically challenged.

Where the evidence is further removed from the original sensor data, additional scrutiny becomes necessary. Courts should then apply the Machine Credibility Test, with the burden resting on the party seeking to rely on the evidence. Any available regulatory documentation, including material comparable to that required under the European Union's AI Act, should also be considered where relevant and a meaningful opportunity to test the evidence should follow. This requires disclosure of technical material and, where appropriate, independent expert examination. Legitimate claims

32 Brandon L Garrett and Cynthia Rudin, *'The Right to a Glass Box: Rethinking the Use of Artificial Intelligence in Criminal Justice'* (2024) 109 Cornell L Rev 561

of confidentiality can be addressed through protective orders, but they should not prevent a party from examining the basis of the evidence relied upon against them.

The court should then determine the weight to be attached to the evidence by applying the Principle of Inferential Distance. High-inference outputs will ordinarily require independent corroboration before they are relied upon to establish guilt. That expectation may be displaced where the Machine Credibility Test demonstrates that the evidence is sufficiently reliable.

Finally, the court should explain its conclusions on every disputed stage of the mediation chain. Clear reasons not only improve the quality of decision-making but also allow meaningful appellate review. They provide an important safeguard against both institutional limitations and the tendency to place excessive confidence in automated systems.

The framework is intended to be introduced gradually. The first stages can be adopted within the existing statutory framework, while later stages would benefit from judicial practice directions, delegated legislation under the Bharatiya Sakshya Adhiniyam, or future legislative amendment. One possibility would be to supplement the existing authenticity certificate with a separate certificate addressing inferential reliability.

XII. Conclusion

The smartwatch worn by Connie Dabate did not witness her death in the way a person could. It recorded physical signals, processed them through software, and generated a conclusion that the prosecution relied upon to argue that she was alive and walking. Each step involved interpretation, and each carried the possibility of error.

Yet the trial focused largely on the overall reliability of the device rather than on the specific processes that produced the disputed output.

Indian evidence law is well equipped to establish the authenticity of electronic records through section 63 of the Bharatiya Sakshya Adhiniyam. Authenticity alone, however, cannot answer every question that machine-generated evidence presents. A record may be genuine even if the inference drawn from it is mistake.

The Doctrine of Epistemic Mediation, together with the Principle of Inferential Distance, the Machine Credibility Test, Algorithmic Cross-Examination, and the Digital Forensic Due Process Framework, is intended to address that unanswered question by providing a structured method for evaluating inferential reliability.

Wearable devices have provided the principal case study because their limitations are well documented and their accuracy often depends on the circumstances in which they are used. The underlying argument, however, extends beyond wearable technology. Smart homes, autonomous vehicles, medical implants, and predictive forensic software are becoming increasingly

important sources of evidence in criminal proceedings. As courts rely more frequently on these technologies, they will need to ask more than whether an electronic record is authentic. They will also need to consider whether the process that transformed physical events into legal evidence was reliable. Meeting that challenge is essential if machine-generated evidence is to support findings of guilt beyond reasonable doubt.