



Zero Witness for lawyers

Every other AI tool asks you to trust a retention policy. This one can run entirely inside your own browser — where **there is nothing to retain, produce or breach.**

Version 2.0 18 September 2026 Coverage figures read live from this deployment

IN ONE SENTENCE

An AI assistant and legal research workspace that can run **entirely inside your own web browser**, so nothing you type is sent anywhere.

WHAT IT ASKS OF YOU

No account. No sign-up, no email, no password — and so no record anywhere with your name on it. The brief checker is free and will stay free.

WHAT IT CHECKS

Citations, quotations, statutes and court rules against **our own copy of the law**. The figures are in §6, and they are read from this deployment rather than typed.

WHO IT IS WRONG FOR

Anyone under a records-retention duty. Nothing is logged, so there is no audit trail and none can be added. See §9.

CONTENTS

- | | |
|--|--|
| § 1 What this product does | § 7 How this differs from ChatGPT |
| § 2 Terms used in this document | § 8 Where this touches the Model Rules |
| § 3 The question underneath all the others | § 9 What it cannot do |
| § 4 The two places the model can run | THE LIMITS |
| § 5 Getting a wrong answer, and what surrounds the model | § 10 The technical detail |
| START HERE | § 11 How to check any of this yourself |
| § 6 The legal work it does | |

Eleven sections, and you are not expected to read them in order. §5 is the one to read if you read only one; §9 is what this tool cannot do.

§ 1 What this product does

Zero Witness is two things on one set of foundations: a private AI assistant, and a workspace for a legal matter that checks its own citations.

In use it looks like any other AI chatbot. You ask a question and an answer comes back — written by a **large language model**, the same kind of thing that sits behind ChatGPT, and which this document calls simply *the model*. You can attach a document, record a call and get a transcript, or hand it a draft brief and have every citation checked.

What is different is underneath. **You choose where the work actually happens** — on a computer we own, like every other AI service, or entirely inside your own browser, on your own laptop, where your words are never transmitted at all. That second option is why this product exists.

What lawyers use it for

- **Checking a brief before filing.** Every citation, statute, regulation and court rule in a draft is checked against a copy of the law we hold ourselves. The draft never leaves your computer.
- **Checking a quotation.** Are the words in your quotation marks the words the court wrote? Answered mechanically — and it flags a perfect quotation that turns out to come from a dissent.
- **Working a matter.** One folder per case: documents, transcripts, conversations, and every authority relied on, each labelled binding or persuasive for your court.
- **Research.** Full-text search across case law, the U.S. Code, the C.F.R., the Federal Register, the federal rules, and the local rules of 241 courts.

What it costs

The brief checker is free, needs no account, and will stay that way. The assistant and the matter workspace are not yet priced — we would rather watch how people actually use them, and hear what is wrong with them, before deciding. If they are charged for it will be a charge for the service itself: this is not funded by advertising and will never be funded by selling what you put into it, which is most of the reason the architecture looks the way it does.

A note on tone. Nothing here is written to impress you. Where a protection is a promise about our conduct rather than a property of the system, it says so in the same sentence as the claim — and §9 lists what this tool cannot do, including the one thing that disqualifies it outright for some practices.

§ 2 Terms used in this document

Six words do most of the work below. Three of them are used loosely in marketing material, and the looseness is where misunderstandings about privacy start.

A model · an LLM · "the chatbot"

Three words for one thing at different distances. The **chatbot** is the window you type into. Behind it is a **large language model, or LLM** — a very large file of numbers plus a short program that multiplies them. This document says "**the model**" for short. It is not a search engine and not a database: it holds no documents and looks nothing up. See §5.

Running a model

Doing that multiplication, which takes real computing power. **The whole privacy question is which computer does it**, because that is the computer your question has to reach.

Open-weight

A model whose numbers its maker has **published for anyone to download and run**. Every model here is open-weight. A closed model can only ever run on its owner's servers.

Browser mode

Running the model **on your own machine, inside the web page**. Downloaded once and kept. After that, your questions and documents are not transmitted — no request carries them.

Server mode

Running it **on hardware we own outright**, never rented cloud. The message is scrambled in your browser before it is sent, and unscrambled only in memory at our end.

Hallucination

A model **producing something fluent, confident and false** — most notoriously a citation with the right shape that refers to nothing. Not a malfunction; §5 explains why.

§ 3 The question underneath all the others

Lawyers rarely ask us whether the model is any good. They ask something narrower: *is putting this document into it a disclosure?*

Every mainstream assistant answers that the same way. It points at a retention policy, a security page, a certification — and asks you to accept that the company will behave as described. That is not a bad answer, and for material that is sensitive rather than privileged it is probably the right trade.

But it has a weakness that matters specifically to lawyers. **A retention policy is a promise about conduct, not a property of the system**. It can be revised, misconfigured, or overridden by a court in litigation you are not party to — as in 2025, when a preservation order in copyright litigation required OpenAI to retain output logs it would otherwise have deleted. That order was contested and later narrowed, and enterprise arrangements were treated differently. The point is not about one company: *a policy is the kind of thing that can change without you*.

This product gives two answers. One is a policy of exactly that kind, stated plainly with its limits admitted. The other is not a policy at all — **there is no request carrying your words, so there is nothing to retain, produce, subpoena or breach**.

§ 4 The two places the model can run

A model has to run somewhere, and where it runs is the whole of the privacy question. There are only two honest options. Most products choose one for you; this one offers both and shows which you are in at all times.

In your browser. A model is downloaded once into your browser's storage and runs on your computer's graphics hardware, inside the page. Ten are offered, from 207 MB to 2.6 GB. Nothing is installed in the

ordinary sense — no program, no icon; clearing your browser's data removes it. Once loaded, your question, your document and the answer never leave the machine. Not "we do not keep them": there is no network request that carries them, and you can disconnect from the internet and keep working. **On hardware we own — and why you would ever choose it.** Three things you actually get. **A far more capable model:** the one in your browser is small enough to run inside a web page, and the interface says plainly what that costs — it extracts and searches well, and it analyses poorly. **Room to think:** a browser model can hold roughly 1,700 words of conversation and document at once, ours around 4,000. **And it works on any machine:** browser mode needs particular graphics hardware and a download before it can answer at all, so on an older laptop, a phone, or a locked-down firm machine it is simply unavailable — the page tells you so rather than failing.

The cost is the obvious one, and it is the whole of the trade: your words leave the device. They are scrambled in your browser before they go, unscrambled only in memory on hardware we own outright, never written to disk, and the model releases what it held when the conversation ends. Specifics in §10.

Which is why a matter has a third setting, and for most legal work it is the one to use. Your documents stay on your machine and are indexed there; when you ask a question, only the handful of passages that actually matched it are sent to the larger model. Neither model can hold a whole case file at once — that is what the index is for, and it works the same way in both modes.

In your browserSTRONGEST	On our hardwareSTRONG
STAYS ON YOUR DEVICE — Everything you type. The question, the conversation, the reply. — Every document you attach. Opened and read by the browser itself. — A matter's documents. Split into passages and indexed on your own machine. — Audio you record or drop in. Transcribed on the device; never uploaded.	WHAT PROTECTS IT — It is scrambled on your own computer, before it travels. Not by the network along the way — by the page itself, so everything in between carries only unreadable text. Each message is locked separately, so opening one would not open another. — Nothing that could unlock it is written down. The secret that undoes the scrambling exists only in our machine's working memory, never on a disk, and is thrown away and remade every hour — so there is nothing to seize, and yesterday's traffic cannot be opened with anything taken today. — We read it to answer, and then it is gone. There is no conversation log — no file to produce in discovery, no backup, and nothing to breach.
CROSSES THE BOUNDARY — A citation, to be checked. The numbers and nothing around them — 786 F.3d 559. Not the sentence it sits in, not the case name, not the document. — The words of a search you run. Scrambled the same way, to our own copy of the law. No record of the query is kept. Those two exceptions are what make citation-checking and research possible at all, and they are named on screen before they happen, every time.	WHAT IT DOES NOT PROTECT AGAINST — Us. We unscramble the message to answer it — a model cannot answer what it cannot read. — This is not end-to-end encryption and we will never call it that. It shuts out everyone in between, including the company that carries the traffic to us and would otherwise see everything in the clear. It does not shut out us. — Proof. "Never written down" describes how we run our infrastructure. It is true, and not provable to you the way browser mode is. In browser mode the guarantee is architectural — it follows from how the thing is built. Here it is operational. Both are real; they are not the same kind of thing.

§ 5 **Getting a wrong answer, and what surrounds the model**

Read this section if you read only one. Privacy is the architecture; this decides whether the tool is safe to rely on.

A large language model is not a library or a search engine. It was built by absorbing an enormous quantity of text and learning, statistically, which words follow which. Asked something, it produces the sequence that best fits the pattern of an answer.

That is why it reads so fluently, and exactly why it invents. **A citation has a very regular shape** — volume, reporter, page, court, year — so a system that has absorbed a million of them can produce a new one with a perfect shape that refers to nothing. It is not lying and it is not broken; from the inside a fabrication is indistinguishable from a real citation, which is why it will defend one if challenged. A larger model makes this rarer, which makes it *harder to catch*, not safer. The sanctions that began with *Mata v. Avianca* in 2023 have recurred steadily since, on newer systems.

So the model is never trusted here. It is fenced. That fencing is the actual product; the model is one component and deliberately the least trusted one. Every factual claim it makes is either checked by something that is not a model, or is visibly marked as unchecked.

The one that looks like a solved problem and is not

A check that only asks "*does something exist at this citation?*" is worse than no check, because invented citations frequently resolve. A real example from this product's own testing, asking for Seventh Circuit authority:

WORKED EXAMPLE — A CITATION THAT CHECKS OUT AND IS STILL WRONG

THE MODEL PRODUCED

Hoffman v. City of Chicago, 642 N.E.2d 1256

THAT CITATION RESOLVES — TO

Illinois State Toll Highway Authority v. American National Bank

A real case, correctly reported at those numbers, with nothing to do with the proposition it was offered for. A checker that stops at "found" marks this **confirmed** — and in doing so **launders the invention**, handing you a green tick on a fabrication.

So the name the draft claims is compared against the name of the case that actually sits there. A mismatch is flagged in red and deliberately carries **no binding-or-persuasive label**, because printing "Binding" beside "wrong name" would assert the authority of a case nobody was talking about.

The layers, in the order they apply

1 Put real law in front of it before it answers

With research on, your question is first run against our own copy of the case law and the actual opinions are placed in front of the model before it writes. This *reduces* invention rather than catching it afterwards — the only layer that does.

2 In a matter, answer only from your own documents

Your files are indexed on your machine; a question retrieves the passages that match and the answer is built from those. The panel beside the answer says which were used, or says plainly that none were.

3 Check every citation mechanically

Not by asking a model whether it is confident — by looking the citation up in a copy of the law on disk. One we cannot find is reported as not found, never as invented, because our coverage is uneven and a real case can be missing.

4 Check the name, not just the numbers

The layer in the example above, and the one most tools do not have.

5 Check quotations against the court's own words

No model involved: the words are in the opinion or they are not, and the passage is printed so you can see the answer rather than take it. A word-perfect quotation from a dissent is flagged as one.

6 Where judgment is unavoidable, label it as judgment

Whether a case supports the sentence it was cited for cannot be answered mechanically. That check does use a model, says so, prints the passages it read, and never words a null result as "this case does not support you". A verdict citing no specific passage is automatically downgraded to "could not tell" — a rule added after a real run produced a false "supported" that was the instruction's own wording read back.

7 Choose the model on measured honesty, not capability

Two candidates were put to six propositions against real opinions, then ten runs on the decisive one. Asked whether a case supported a proposition appearing nowhere in it, one answered **supported in 4 runs of 10** — inventing the passage's contents to match. The other did so **0 times in 10**. The legal workspace defaults to the second *despite it scoring lower elsewhere*, because here the only unrecoverable failure is a confident wrong answer.

8 Read the document for text you cannot see

A brief you were sent can carry instructions aimed at whatever AI reads it next — white on white, characters with no width, a rendering mode that draws nothing. Invisible to you, present in the file, and it would otherwise reach the model. Looked for in your browser, with nothing sent to do it, and the warning accuses nobody: hidden text is also how templates carry instructions and drafting notes get left behind.

None of this makes the model reliable. It makes its **factual claims checkable** and its **unchecked claims visibly unchecked** — every judgment is printed with the text it was made from, because the design assumes you will check the work rather than accept it. Its reasoning is verified by nothing, and neither is anything it says about a document it was shown. In your browser it is a small model and a correspondingly weak analyst. Treat the output as a memo from a capable first-week associate who occasionally invents a case with total conviction: useful, and signed by you.

The brief checker is free and needs no account. Drop in a brief or paste it; every citation is found in your browser, and you are shown the exact list that will be sent before anything is sent — only ever volume, reporter and page. The document never leaves the tab. It covers case citations, the U.S. Code and C.F.R., the Federal Register, the federal rules, and local and state court rules from 241 courts.

A matter workspace holds the documents, transcripts, conversations and authorities for one case, each authority labelled binding or persuasive for the court you named. Each matter carries a posture you set: *on this device* (documents and reasoning both stay in the browser; only a citation and a search leave), *documents here, thinking there* (files stay and are indexed locally, a far more capable model answers and sees only matched passages), or *kept on our hardware* — declared and deliberately not built, since it would make us custodians of privileged material.

What our copy of the law holds

Everything resolves against data on our own disk before anyone else is asked. These figures are read from this deployment as the page loads, so they cannot drift from what a check can actually find. **Every court, code and rule set — and every court we do *not* hold, with the reason — is listed on [the coverage page](#).**

SOURCE	HELD	NOTES
Case law	8,872,601 citations	145 courts · quarterly snapshot
Opinions with full text	3,929,256	Searchable word by word
Federal opinions by docket	112,685	From GovInfo, read daily
U.S. Code and C.F.R.	287,282 sections	A dated snapshot, not the live text
Federal Register	328,792 documents	From 2015
Federal rules	544 rules	Civil, criminal, appellate, bankruptcy, evidence
Court rules	41,196 rules	241 courts, from each court's own site

Coverage is uneven and the page says where. Federal district decisions in the bulk snapshot stop between 2017 and 2019 — a cliff, not a taper — which is why recent opinions are read daily from separate sources. **A citation we cannot find may be perfectly real**, and the tool says so rather than implying a fabrication.

§ 7 How this differs from ChatGPT

Worth making this comparison fairly. Mainstream assistants are enormously more capable than anything that fits in a browser tab, and their enterprise tiers are serious offerings with real contractual protections. The difference is not that they are careless — it is *where the work happens*, and therefore what kind of assurance is available at all.

	ZERO WITNESS — BROWSER	ZERO WITNESS — OUR HARDWARE	CONSUMER ASSISTANTS	ENTERPRISE / TEAM TIERS
Where the model runs	Your own computer	Hardware we own; no rented cloud	The vendor's cloud	The vendor's cloud
Are your words transmitted?	No. No request carries them	Yes — encrypted in your browser first	Yes	Yes
Retention of your inputs	Nothing to retain	Nothing written to disk; no log exists	Retained by default; deleted on a schedule	Contractual; zero-retention options common
Training on your inputs	Impossible — they never arrive	No. No pipeline and no corpus	Commonly yes by default; opt-out offered	Contractually no, by default
Reachable by a third party's subpoena or preservation order	Nothing to reach	Nothing stored to produce	Yes — logs exist and have been ordered preserved	Narrowed by contract; logs may still exist
Account required	No	No	Yes	Yes, plus firm administration
Model capability	Small. Extracts and searches well; analyses poorly	A 27-billion-parameter open model	Frontier-scale; far more capable	Frontier-scale
Citations and quotations checked against real law	Yes — against our own copy	Yes	No; fabrication is the known failure	Varies; usually no
Audit trail for records retention	None, and none can be added	None, and none can be added	Some	Yes — a real reason to prefer them
SOC 2 / formal certification	None	None	Varies	Yes
Cost	Brief checker free; the rest not yet priced	Not yet priced	Free and paid tiers	Per seat

The honest summary. If your obligation is to keep a record of what was done, an enterprise tier is the better tool and this one cannot be made into it. If your obligation is to keep privileged material from leaving the building, browser mode does something no cloud product can do at any price — and the last three rows are what you pay for it.

§ 8 Where this touches the Model Rules

This is not legal advice and not an opinion about your obligations. We are not your lawyers and cannot be. What follows is a plain statement of facts about this tool bearing on the rules lawyers have raised with us, so you can apply them yourself. Opinion 512 is the national anchor and the ABA's Task Force on Law and Artificial Intelligence has published further reports since (most recently December 2025), but **the operative rules are your state's** — more than thirty-five bars have now issued AI guidance, and it varies.

Rule 1.6(a), (c)

Confidentiality; reasonable efforts

Rule 1.6(c) asks for **reasonable efforts**, not a guarantee that nothing ever leaves the office. Comment [18] lists the factors: sensitivity, likelihood of disclosure absent safeguards, their cost and difficulty, and whether they impede representing the client. Browser mode is unusual against that list — the safeguard needs no configuration and no third party, and the likelihood of disclosure in transit is not reduced but absent, because there is no transit.

Formal Opinion 512

ABA, July 2024 — generative AI tools

The ABA's first ethics guidance on generative AI. On confidentiality it asks a lawyer to keep confidential everything relating to a representation *regardless of its source*, and to obtain the client's **informed consent before using a third-party generative AI program where that would disclose confidential information**. That conditional is where this tool sits: in browser mode nothing is disclosed to a third-party program, so the trigger is not reached on that limb; on our hardware it plainly is, and §10 states exactly what is retained and what is trained on — nothing, in both cases. Which mode you are in is shown at all times, and it is your choice rather than our default.

Rule 1.1, Comment [8]

Competence in technology

The facts a competent evaluation needs are in §4, §5 and §10 rather than scattered across a privacy policy, a security page and a sub-processor list. **The coverage page names every court and code we hold and every one we do not**, with dates, so the limits of the research are inspectable rather than asserted.

Rule 1.4

Communication; informed consent

Whether a use requires telling the client, or consent, is your judgment. What this gives you is a **precise answer to give them**: in browser mode, that their material was not transmitted; on our hardware, that it was encrypted in transit, read in memory, never written down, and not used for training.

Rules 5.1 & 5.3

Supervision of nonlawyer assistance

Rule 5.3 reaches outside vendors. The problem is smaller here in one respect and unchanged in another: there is **no vendor-side store of client material to govern**, but the output needs the same review any junior work product would — see §5 on which parts of an answer are checked and which are not.

Rule 3.3 & Fed. R. Civ. P. 11

Candor; certification of filings

The fabricated-citation sanctions beginning with *Mata v. Avianca* are what §5 exists to address. **A check is not a substitute for reading the case**. These tools tell you a case exists, which court decided it, whether the claimed name matches, and whether your quoted words are in it. They cannot tell you it remains good law, and they do not sign your filing.

Rule 1.5

Fees

Opinion 512 addresses billing where AI saves time: you may charge for the time spent putting a matter into the tool and for reviewing what comes back, but **in most circumstances you cannot bill a client for learning to use it**. Nothing here changes that. We note only that the brief checker is free and needs no account, and the workspace is not yet priced — if it is charged for it will be a charge for the service, never funded by advertising or by selling what you put into it.

Standing orders

Judge-specific AI disclosure

Many judges now require disclosure or certification when generative AI is used in a filing, and the orders vary widely. **We do not hold them and do not check them for you**. That remains yours to look up.

Illinois

Supreme Court AI Policy, Jan 2025 · ARDC guide, Oct 2025

The Illinois Supreme Court's policy says lawyers' use of AI **"may be expected, should not be discouraged, and is authorized"** where it meets legal and ethical standards, and that **disclosure of AI use should not be required in a pleading** — a markedly different line from the federal judges who require certification. Accountability for the final work product is untouched, and output must be reviewed before it is filed. The ARDC's *Illinois Attorney's Guide to Implementing AI* then gives the criteria for choosing a tool: **model-training settings, data retention, isolation and vendor terms**, and whether it is

managed by a third party or hosted internally. §10 answers that list line by line, and browser mode is the far end of "hosted internally" — not the firm's own server but the lawyer's own machine, with no vendor receiving anything. There is no ISBA ethics opinion on AI; its standing committee publishes FAQs and best practices instead.

§ 9 What it cannot do

Read this before relying on any of it. The fastest way to lose a careful reader is to let them find a gap themselves.

It is not a citator

It says a case exists and which court decided it — not whether it is still good law. No Shepard's signal, no KeyCite, no treatment history: **an overruled case looks exactly like one that is not**. No free source carries that, and we will not build a partial version, because a partial treatment signal is more dangerous than none.

There is no audit trail, and none can be added

Nothing is logged and nothing kept on our hardware. If you are under a records-retention duty — many regulated practices are — **this is the wrong tool and cannot be made into the right one**. That follows from the architecture; it is not an unbuilt feature.

Coverage is a snapshot, and uneven

Recent decisions, unpublished opinions and state trial courts are thin, and not every state's court rules are held. A citation the tool cannot find may be perfectly real, and the screen says so.

Binding and persuasive are applied mechanically

A federal district court is bound by its own circuit; a state court is bound by no federal court except the Supreme Court on federal questions. That is all the labels mean — **they take no account of whether a case is on point**.

The model gets things wrong, and §5 is only a fence

Citations and quotations are checked; **its reasoning is not**, and neither is anything it says about a document it was shown. In your browser it is a small model and a correspondingly weak analyst.

No certification, insurance or negotiated contract

No SOC 2 report, no BAA, no data-processing agreement today. If your client's outside-counsel guidelines require one, this does not satisfy them — cheaper to say here than to discover in the fourth meeting.

It does not file, docket or calendar

No PACER, no ECF, no deadlines, no conflicts check.

§ 10 The technical detail

For a reader who wants specifics, or an IT adviser reviewing this on your behalf. Nothing here changes the conclusions above.

Encryption between your browser and our hardware

The page and the server perform an **ECDH** key agreement on the **P-256** curve, derive two directional keys through **HKDF-SHA256**, and encrypt with **AES-256-GCM** — one key for the request and one for the reply, so a request key cannot forge a reply. Both public keys are bound into the derivation, so the keys belong to that one pairing and cannot be replayed. The browser generates a **fresh keypair for every message**; ours exists in memory only, is never written to disk, and is replaced hourly with one period of overlap so a page left open does not fail mid-conversation. There is deliberately **no unencrypted fallback path** — a downgrade option is a downgrade attack.

What this buys is specific. Traffic reaches us through a content-delivery network, which terminates transport-layer encryption at its own edge, so without this every message would be readable in the clear at a third party we do not control. This closes that, and nothing else. It is **not end-to-end encryption**, because we decrypt to answer.

Storage on your device

Conversations, transcripts, matters and indexed documents are written to your browser's own storage, on your machine, synchronised nowhere. Clearing site data removes them; we never hold a copy. The store can optionally be encrypted at rest under a passphrase you choose: `PBKDF2-SHA256` at **600,000 iterations** (the current OWASP figure) derives a key wrapping a per-device data key, and every record is sealed with `AES-256-GCM`, bound to its own location so a row cannot be moved. A recovery key is shown once; losing both it and the passphrase means the data is unrecoverable, *including by us*.

Precisely: it protects against someone holding your computer, or a copy of its files, without the passphrase. It protects nothing while unlocked, does not replace full-disk encryption, and we do not claim the browser has discarded every earlier unencrypted copy — storage engines compact on their own schedule. You can also export everything to one encrypted file, or mirror it to a folder on your own disk.

What exists on our side

- **No conversation logging, ever.** Not truncated, not hashed, not sampled. **No search-query logging** either — a search term is conversation content by another name.
- **No training.** Published open-weight models, run unchanged. There is no fine-tuning pipeline and no corpus of conversations that could feed one.
- **An exception log** carrying an exception type, a fixed operation name and a stack trace from our own machine — never a message, search term, attachment or address. Email and IP scrubbing is a backstop, not the mechanism: what keeps content out is that no call site passes any.
- **One operational count:** how many conversations queued for hardware, tallied once per calendar day. No time of day, no address, nothing per request.

Analytics, and on-device processing

Page views are counted with Plausible — no cookies, no cross-site profile. The tracker is served from our own origin and its events relayed by our own server, so *your browser never contacts theirs*; they still receive the visit and the address it came from. Alongside the view we count how long the page was in use and how far it was scrolled. Nothing you *do* is counted: not a matter's name, a question, a document, a citation, or which part of the workspace you opened. It is the one item here you can switch off at your end, with any ordinary blocker.

Transcription runs on your machine using Whisper; audio is never uploaded. A matter's documents are indexed locally with a 23 MB embedding model and matched locally, so in browser mode not even the matched passages leave. Word, PDF, spreadsheet and text files are all parsed in the browser.

§ 11 How to check any of this yourself

You should not take a privacy claim on the strength of a document describing it. Browser mode is unusual in that the central claim is **verifiable from your own machine in about two minutes**, with tools already installed.

1 Open the network inspector and count the requests

Right-click, Inspect, Network tab. Load [the chat page](#), switch to "In your browser", let the model download, then ask several questions. Nothing appears carrying your text — a whole conversation's only request is for an encryption key.

2 Turn off your Wi-Fi and keep talking

The most direct test there is. With the model loaded, disconnect the machine entirely and carry on. It keeps answering, because it was never asking anyone.

3 Read the list before it is sent

On [the brief checker](#), drop in a real brief. Before anything leaves, the page shows the *exact* list of citations that will go, as numbers. Compare it against the document — that screen is the product.

4 Read the coverage page, court by court

[Every court, code and rule set](#) we hold, with counts and dates — and every one we do not, with the reason. It is generated from the same files the tool answers from, so the figures cannot drift from what a check can find.

5 Read the measurements on the models

The browser models are [benchmarked and published](#), including the bad results. One states recent news as fact in two runs of five, and its tile says so where you would choose it.

If anything here does not match what you observe, we want to hear about it before you do anything else with it. That is not politeness — a claim that survives being tested is the only asset this product has.

No account. No training. No witness. Check a brief Legal coverage Questions