

ENI6MA : How Naive Attack Intuition is Wrong

Dr. Lin Wang, Dr. Dawn Lipscomb, Dr. Anish Mohamad,
Dr. Francisco Perez, Frank Dylan Rosario

September 2025

*For challengers who worry that (1) an attacker can learn the password and the private map by watching multiple successful logins, or (2) a determined attacker can randomly guess both password and private hologram. We describe why neither path works, **not mathematically, not statistically, not computationally, and not physically.***

Introduction

Most people carry around a mental picture of authentication that comes from passwords and PINs: you type a fixed string, a computer checks it, and you're in. If a camera, keylogger, or phishing page captures that string, game over. From within that picture, a natural fear follows: **"If an attacker watches enough successful logins, they'll eventually learn the secret, or they can just keep guessing until they hit it."** This essay tackles those worries for the ENI6MA rotating-ring method and explains, in minimal mathematics and plain language, why both intuitions are wrong, not just in theory, but in practice and even in principle.

ENI6MA asks you to prove knowledge of a secret **without exporting the secret**. Imagine a single alphabet arranged on a ring, letters around a clock. Each login consists of a few rounds. At the start of every round, the system **rotates** the ring by a fresh, unpredictable amount and paints the circle into **C zones** (like colored slices). Your secret is a sequence of letters. To advance a round, you indicate **which zone** currently contains the next letter. Crucially, you do not say the zone number out loud; you emit a **witness token** (a word, icon, or gesture) drawn from a set only you understand because you keep a **private map** that ties tokens to zone indices. Outsiders can watch the layout and hear the token, but they **don't know** how that token translates to a zone, and the layout changes again in the very next round.

Two design principles follow. First, **ephemerality**: because rotations are fresh and time-tied, yesterday's geometry won't return; a recording of yesterday's success is a souvenir, not a skeleton key. Second, **masking**: witnesses

are public labels without a legend; even perfect notes about what was said do not reveal which zone was meant. Those two facts already weaken the everyday intuition that “watching more reveals more.” $\hat{M}$ = Private Map

We address two naive objections and challenges. The first addresses **Observation**: the claim that, with enough recorded sessions, an attacker can reconstruct both the password and the private map. We’ll show why the observable data, layouts and tokens, have the **same statistical behavior** no matter which secret you used. When the distribution of what you can see doesn’t depend on the secret, no amount of watching creates a signal where none exists. Frequency tricks don’t bite because rotations spread all letters evenly across zones over time; correlation tricks don’t bite because each round is independent; alignment tricks don’t bite because the private map never leaks and the ring keeps moving. Even sophisticated machine learning only rediscovers that the layout generator is random; it does not uncover your secret.

The second argument addresses **Guessing**: the idea that a determined attacker can brute-force the password and the map. Here the protocol is engineered to be hostile to luck. The space of possibilities (secrets of length L over N letters, times $C!$ private maps) grows explosively, while the chance of a blind guess passing L rounds is roughly $(1/C)^L$, tiny by design. Add in cryptographic randomness that drives the rotations, and the search balloons beyond what computers, even fantasy-level ones, could attempt in the lifetime and energy budget of a civilization. Rate limits then turn “astronomically unlikely” into “practically impossible.”

Along the way, the essay introduces a few gentle ideas, **uniform rotation** (every letter spends equal time in every zone), **independence** (each round’s geometry is fresh), and **unidentifiability** (many different pairs of secret and private map can “explain” the same recording). You won’t need heavy algebra; the goal is clarity, not technical bravado. We’ll lean on simple analogies: a dance that must match **tonight’s** song; a museum rehung overnight; six unlabeled buttons that make six sounds only you have mapped to numbers.

The thesis is simple: ENI6MA turns a secret from a **possession** (a reusable string that leaks on contact) into a **performance** (a short, per-session demonstration in a changing world). Recordings of past performances don’t open future doors. Guessing the right performance is both improbable by construction and infeasible in our universe. That is why the two worries, learning by watching and winning by guessing, fail here: not mathematically, not statistically, not computationally, and not physically.

“Why Recording Valid Proofs Won’t Help”: Observation and Guessing Both Fail Against ENI6MA’s Rotating-Ring Authentication Cypher.

A quick, plain-language primer (so we’re talking about the same thing)

Imagine the alphabet printed once, in a fixed circular order, like letters around a clock. Each login is a short ceremony split into rounds. At the start of **each round**, the verifier (the system) secretly **rotates** the ring by a fresh, unpredictable amount and then paints the circle into **C zones** (think colored slices). Your secret is a sequence of letters. To prove you know it, in each round you simply indicate **which zone** currently contains the next letter of your secret.

There is one more ingredient: a **private map** between zone numbers and the **witness tokens** you actually present (spoken words, button icons, gestures). Outsiders can hear or see the witness tokens you emit, but they **don’t know** how those tokens correspond to the zones they’re watching on the screen or tabletop. Said differently, your “UP” might mean zone 3 to you, but nobody else knows that mapping.

Two design consequences are crucial:

1. **Ephemerality.** The layout (which letters are in which zones) changes every round, and it’s regenerated from fresh randomness tied to time. Yesterday’s successful path is useless today.
2. **Masking.** Outsiders observe only tokens (“UP”, “LEFT”, etc.), not the underlying zone indices. Without your private token→zone legend, the tokens are just unlabeled categories.

With that picture in mind, let’s engage the two worries head-on.

Objection [1]: “If an attacker records enough valid logins, they’ll reconstruct the password and the private map.”

This sounds plausible if you imagine traditional passwords, where watching keystrokes really does help. Here, it doesn’t. The reason is simple but deep: **what the attacker can record has a probability distribution that does**

not depend on your secret. When the observable data have the *same* distribution no matter which secret you chose, no amount of watching creates signal out of noise.

Let's unpack that in everyday terms.

What the camera actually captures

A camera (or a person taking notes) can record two things per round:

- The **layout**, which letters fell into which colored zone *this round* after the system's rotation.
- Your **witness token**, the public label you emitted (e.g., "UP"), which is secretly tied, via your private map, to the zone you meant.

They can also note whether the overall attempt passed or failed.

That's it. They do **not** learn your private token→zone map, and they cannot reuse the recorded performance later because the next session's geometry will be freshly regenerated.

Why statistics don't accumulate signal

The core trick is the **rotation**. If each round's rotation is unpredictable and independent of previous rounds (which the system ensures), then **every letter spends equal time in each zone across rounds**. If there are six zones, your letter is in each zone about one-sixth of the time. No letter prefers any zone, and the *sequence* of zones a particular letter visits is just a string of fair, independent draws.

When you speak a witness token, you're just relabeling that fair draw via your private map. Outsiders don't know the map, so the tokens they hear are also just fair, independent draws over the token set. It doesn't matter whether your next secret letter is "E" or "Z", from the outside, the statistics of what you say look the same.

Put starkly: **an observer's data look identical no matter which secret you used**. If every secret produces the same observable pattern, **no amount of observing distinguishes secrets**. Collecting more videos grows the pile of data but not the information about your secret.

“But I’ll do frequency analysis!”

Classic frequency attacks count how often symbols appear. That’s how code-breakers attack simple substitution ciphers. Here, frequency analysis has no bite:

- Because the ring is randomly rotated each round, each letter lands in each zone equally often.
- Because outsiders don’t know your private map, they cannot align a token (“UP”) with a specific zone (“red”) across rounds.
- Result: token frequencies are flat and uninformative. Every secret induces the same flat histogram.

“What about correlations, like bigrams (‘T’ often followed by ‘H’)?”

Correlations don’t survive either. The system rotates the ring independently each round, so the zone for the next letter is independent of the zone for the previous letter. Even repeated letters (like the double “O” in “NOODLE”) won’t show up as a pattern: each “O” occurs under a different rotation, so their zones (and thus their tokens) are independent draws again. The famous fingerprints that cryptanalysts exploit simply don’t appear.

“Okay, but I see the layouts *and* the tokens, can’t I align them over time?”

To align tokens to zones, you need a stable anchor: some way to say “this token always meant that zone.” The protocol denies you that anchor in two ways:

- **Private map.** The mapping from tokens to zone numbers is secret and never leaves the user. You’re hearing labels without a legend.
- **Fresh rotations.** Even if you temporarily guess a token’s meaning in one round, the next rotation moves **every** letter in lockstep to new zones. Your guess is immediately scrambled.

Even clever cross-round stitching doesn’t stick. Each round is a fresh coin flip from your vantage point. There is **nothing** to accumulate.

The identifiability trap (why many “explanations” fit the same video)

There’s a deeper snag for the eavesdropper: many different pairs of (*secret*, *private-map*) can “explain” the same observed tokens after the fact. Swap two

token labels in your private map, and you can swap the corresponding letters in the hypothesized secret to match the same transcript. This built-in **symmetry** means the outside observer can never pin down a unique answer, even if they entertained wild brute-force reconstructions. They end up with a cloud of equally plausible stories, all related by relabeling. In statistical terms, the model is **unidentifiable** from passive data.

“But I have machine learning models!”

A perfect classifier can’t extract information that isn’t in the inputs. If the layout and token sequences are statistically independent of your secret (and that’s what the protocol ensures), then any learning system will simply rediscover properties of the layout generator, not the user’s secret. More data just sharpens the estimate of a distribution that *doesn’t depend on the secret*. You don’t beat a zero-information channel with more samples.

“What if I record for a very, very long time?”

You’ll confirm that tokens look uniformly random and that the system keeps rotating the ring freshly each round. That’s it. When the expected pattern is flat, the law of large numbers only gives you a **flatter** flat line. Watching forever does not conjure structure out of noise.

If you’re worrying about a tiny imperfection, like some zones occasionally having one extra character when the alphabet doesn’t split evenly, two notes calm that fear:

- The bias is microscopic (at most “one extra out of the whole alphabet”), and **you don’t know which token corresponds to which zone anyway**.
- If anyone still frets, implementations can **pad** the alphabet (add a few dummy symbols) so zones are exactly the same size. The microscopic wrinkle disappears entirely.

The replay myth

A cousin of the “watch more” claim is “just replay the successful session.” That fails by design. A replay would need the next session’s ring to rotate by the **same** secret amounts as last time. Because rotations are derived from fresh randomness tied to time, the chance of an accidental exact repeat is so small that, in practice, it never happens. Yesterday’s performance is a fossil: interesting to watch, **inert** against today’s door.

Bottom line for Worry [1]

The **observable** world, layouts and tokens, has the **same** statistics for every secret. Without a difference in distribution, there is **nothing to learn**, no matter

how long you watch. That’s the end of the observation attack: **not mathematically** (no dependence), **not statistically** (no signal), **not computationally** (no algorithm can compute what isn’t there), and **not physically** (no lab, no camera, no GPU changes the fact that the channel carries zero information about the secret).

Objection [2]: “Even if observation doesn’t help, can’t a determined attacker just guess the password and the private map?”

Now we switch from “learning by watching” to “forcing with brute force.” Two levers exist: guess the secret; guess the private map. We’ll show that the combined space is explosively large, the per-attempt success probability is intentionally microscopic, and the resource demands to “just try them all” blow past the limits of computation, energy, and time.

How big is the guessing space?

Even without equations, the sizes balloon fast:

- If your secret is a sequence of length L over N letters, there are N^L possible secrets. That grows exponentially with length.
- The private map from **C zones** to **C tokens** is a **permutation**, which means there are $C!$ (C factorial) different maps. Even for modest C (say 6 or 8), that’s hundreds to thousands of possibilities **for the map alone**.

To guess **both**, an attacker faces the product: “secrets $\times$ maps.” Take a friendly example: letters $N = 26$, secret length $L = 8$, zones $C = 6$. The total is in the trillions. Increase the length to 10 or 12, or raise C to 8 or 10, and the count rockets into numbers that simply can’t be enumerated in any sensible time.

And remember the identifiability trap from above: multiple pairs (secret, map) can explain the **same** observed history. Even if someone tried to rule out candidates by comparing to recordings, they’d be left with a huge herd of equally plausible survivors rather than a single winner. Enumeration doesn’t converge to “the truth”; it converges to a **swamp**.

The per-attempt success chance is engineered to be tiny

Every round, a blind guesser picks a token hoping it matches the zone containing your next letter. With C zones, the chance per round is 1 out of C . Over L rounds, the chance you get **every** round right is about $(1/C)^L$. For six zones across six rounds, that’s roughly one in a **million**. And that’s before basic rate-limiting (“you get only a few tries per hour”) or lockouts.

The protocol also uses cryptographic randomness (like a large secret seed) to generate the rotations, which adds another protective layer. Even if someone somehow guessed all L tokens correctly, they would still need their attempt to match the cryptographic state that defines the exact geometry; the chance of stumbling into that match is effectively zero. You don't "luck into" a 256-bit secret by guesswork.

“But computers are fast! We’ll parallelize.”

Let’s give the attacker more credit than they deserve and do a back-of-the-envelope calculation.

- A cutting-edge supercomputer can do around a **quintillion** operations per second (that’s 10^{18}). Run it flat-out, 24/7, for 30 years ($\sim 10^9$ seconds), and you get $\sim 10^{27}$ operations.
- To search a 256-bit cryptographic space (the seed that drives rotations), you’d need on the order of 2^{256} tries on average, about 10^{77} . That’s **fifty** orders of magnitude more than your whole 30-year exascale run can deliver.
- “Okay, so we’ll use a billion such supercomputers.” You’d still be short by **forty** orders of magnitude.

That’s before you pay for storage, networking, cooling, and the small matter of having a billion exascale machines. Which you don’t.

Energy, heat, and the laws of physics

Even if somebody found the money for a planetary-scale data center, **energy** kills the plan. Real computers erase bits (a lot of them). Erasing information produces heat, there’s a minimum energy cost per erased bit baked into physics. Comb through numbers like 10^{77} attempts, and you’re staring at energy budgets so outrageous that, figuratively (and depending on assumptions), you’d **boil oceans** long before you finish. This isn’t hyperbole; it’s the unavoidable bill of doing that much irreversible computation. The planet becomes the bottleneck.

“What about quantum computers?”

A fair question. Quantum search (Grover’s algorithm) can quadratically speed up guessing for some problems. Quadratic sounds big, but it’s not magic. Halving the exponent of an astronomical number still leaves an astronomical number. If your cryptographic seed is 256 bits, Grover’s makes it “like” 128 bits, still way out of reach with realistic error-corrected quantum machines. Protocols simply choose parameters (e.g., 256-bit security and adequate rounds) that remain comfortable against such speedups. Meanwhile, the observation channel still carries zero information, quantum or not.

“Maybe I don’t need the seed; I’ll guess only the secret and the map.”

Even that is rough. With secrets of length 8–12 and moderate alphabets, the number of possibilities is already massive, and the per-attempt acceptance chance $(1/C)^L$ keeps you locked out almost always. Rate-limiting turns “probability of rare luck” into “effectively never in anyone’s lifetime.” You might get a fluke once in tens of millions of tries; you won’t be allowed tens of millions of tries.

The human factors that do matter, and are controlled

The only practical way to weaken guess resistance is to mess with the **parameters**:

- Make L too small (too few rounds).
- Make C too small (too few zones).
- Let people reuse the same session geometry (no fresh rotations).
- Leak the private map by accidentally labeling tokens with on-screen colors or positions.

Competent implementations avoid these pitfalls. Choose reasonable parameters, keep rotations fresh, keep the token→zone map private, and the guess attack becomes a non-starter both in math and in the real world.

A sanity check: “Could someone ever get lucky?”

Luck exists, and lotteries sell tickets for a reason. But the protocol is designed so the “jackpot” odds per attempt are so low that, with basic throttling, **nobody** sees a win in practice. You’re more likely to see hardware failures, power outages, or policy lockouts than a successful blind forge.

Bottom line for Worry [2]

Brute force requires exploring an explosive space (secrets $\times$ maps $\times$ cryptographic states) while each login attempt has a tiny success probability and the system throttles attempts. The needed computation dwarfs our machines; the energy and heat dwarf our planet’s patience. **Not computationally** (too big), **not physically** (too hot, too long), **not statistically** (per-attempt odds stay tiny), and, for good measure, **not mathematically** (the protocol’s construction sets these odds by design).

The deeper reasons both worries miss the mark

It helps to name the conceptual pivots that make this style of authentication behave so differently from passwords.

From “possession” to “performance”

A password is an object you can copy, see it once and you own it. Here, identity is a **performance**: demonstrating that you can navigate *today’s* geometry with *your* private legend. Performances don’t replay. A recording isn’t a key; it’s a souvenir of a show that has ended.

Global motion: a single change moves everything

Because the ring rotates as a whole, a tiny change (one more “tick” of rotation) **moves every letter**. That “rigid body” motion smears out any local pattern you might hope to track, there’s no single letter that stands still long enough for you to correlate against it. Any small bias you thought you found would have to move *with the whole ring*; it doesn’t. That’s why frequency and correlation attacks find no foothold.

The private legend that never leaks

You can’t turn what you hear (“UP”, “DOWN”, “LEFT”, ...) into zone numbers without the private legend. The protocol never shows it, never stores it server-side, and never encodes it in any UI metadata an observer could latch onto. Without that mapping, tokens are just unlabeled categories; interchangeable, unalignable, unhelpful.

Anti-replay “inside” the math, not bolted on

Plenty of systems try to block replay by adding fences (timestamps, IP checks, limits). This design builds replay resistance into the **geometry**. Change the time → change the rotation → change the manifold → invalidate yesterday’s path. The attacker never gets a second use out of a first success.

Auditability without risk

You might worry that if recordings are useless to attackers, they’ll be useless to auditors. The opposite is true. The verifier can keep **receipts** (the non-secret parameters that define each session’s geometry), enough for an auditor to check that every witness token really matched a zone containing the right letter, without keeping any reusable secret, and without knowing the private map. That’s “trust, but verify,” without creating honeypots.

Active attacks, and why they don’t rescue the two worries

So far we discussed passive observation and brute force. What if an attacker tries to **relay** traffic (a phishing site) or stand in the middle?

- A relay still faces the fact that the **real verifier** uses its own fresh rotations tied to its own clock and seed. The imposter cannot keep the fake and the real geometry synchronized across rounds.
- There are variants where the verifier also emits a witness (a “mutual dance”). A middleman then has to fake **two** private maps consistently. That’s not doable without knowing private maps they were never shown.

These are engineering choices, not magic. But they show that even “fancier” attacks are still blocked by the same bones: freshness, private legends, and geometry that the attacker cannot control or predict.

Practical guidance that keeps the guarantees true

If you’re designing or reviewing such a system, here’s the short, non-mathematical checklist:

- **Freshness:** derive each round’s rotation from a high-entropy per-session seed mixed with trusted time. No repeats, no correlations.
- **Separations:** keep the private token→zone mapping in the user’s head/app only. Never mirror that mapping in server logs, URLs, CSS classes, color names, etc.
- **Parameters:** pick enough zones and rounds to keep the per-attempt success probability tiny (e.g., six zones across six rounds already sends it plummeting).
- **Padding:** if you want the cleanest theory, pad the alphabet so it divides evenly into zones. (This is a nicety; it just removes a tiny, irrelevant wrinkle.)
- **Hygiene:** avoid side channels, consistent UI timing, don’t encode zone identity in audio frequencies or pixel artifacts an adversary could analyze.
- **Throttling:** rate-limit attempts and lock out rampant failure. It turns “rare luck” into “practically impossible.”

Do these, and the two worries remain solved in practice because they’re solved in principle.

Thought experiments that make the intuition stick

The dance-floor doorman

The bouncer lets you in only if you perform a five-step dance **to tonight’s song**. You and the bouncer share a private legend (“on ‘UP’ I mean step 3,” etc.). The DJ picks a new song, and with it, a new beat, every minute. Someone can film

you dancing, but tomorrow’s song has a different beat. A replay is off-beat and rejected. Without your private legend, an outsider can’t tell which call (“UP!”) matched which step.

The museum that rehung overnight

You point to the room that contains the next painting in your personal list. Overnight, the curator reshuffles the entire gallery by a rule only they know. The next day, pointing to yesterday’s room number doesn’t find the right painting. Filming your path yesterday doesn’t grant access today.

The label-less buttons

Six buttons emit six sounds. You and the verifier agree privately which sound means which number. Outsiders can record sounds forever; without your private legend, those sounds are interchangeable. Tomorrow the machine rewires which paintings are behind which numbered doors, but the sounds don’t change, and outsiders still don’t know which sound stands for which number.

Each story carries the same moral: **fresh geometry + private legend $\Rightarrow$ recordings and guesses don’t help.**

The four “objections” addressed

When someone says, “If I watch enough, I’ll learn it,” or “If I guess enough, I’ll get it,” we point them to these four layers of impossibility:

1. **Not mathematically.** The observable data (layouts, tokens, pass/fail) are generated by a process whose distribution does not depend on the secret. If the distribution is the same for all secrets, mathematics says there is nothing to infer.
2. **Not statistically.** Frequencies, correlations, and cross-round patterns are intentionally flattened by independent rotations and masked by a private legend. Longer observation only tightens the flatness.
3. **Not computationally.** The combined guessing space (secrets $\times$ maps $\times$ cryptographic states) grows explosively; the per-attempt success odds are tiny; and the system rate-limits tries. No realistic computer farm can churn through it.
4. **Not physically.** Even absurdly optimistic computations would demand energy and time far beyond what a planet or a civilization can supply. Heat, power, and lifetime limits stop the fantasy cold.

Closing: The right mental model

This style of authentication is not a new kind of password; it’s a new kind of **proof**. It asks: *Can you navigate today’s rotating map using your private legend, round after round, without slipping?* Because the map rotates anew and the legend never leaks, a bystander cannot infer the legend by watching, and a brute-forcer cannot stumble into the right path before the universe sends them the electricity bill.

So when a naive (but reasonable!) person says, “Surely, if I watch enough, I’ll get it,” the kindest answer is:

You can watch a million dances, but the song keeps changing and the steps are called in a language you don’t know. And even if you try to fake it, the bouncer keeps changing the music, and you only get a few shots before they stop the show.

That’s why both objections, neither learning by observation and winning by guessing, are defeated here. Not by obscurity, not by wishful thinking, but by **construction**.

Why Naive Intuition About Attacks on ENI6MA are Wrong.

In most systems, a secret (a password, seed phrase, or private key) is something you *show*, even if indirectly, so a camera or keylogger can copy it. If you capture enough keystrokes or screenshots, you’ll piece together the secret. That intuition is **false** here because this protocol never exports the secret itself. Instead, it choreographs a short interaction that proves the user knows the secret *without revealing it*. And critically, **every session takes place in a newly randomized geometry**, the “stage” changes each time, so an old performance cannot be replayed.

This is the central theme: **your knowledge is manifested as a path through a changing landscape, not as a reusable string**. A video of yesterday’s path in yesterday’s landscape doesn’t translate to today.

To make that precise, we need a modest amount of notation and a few simple facts.

The setting (gentle formalism)

We’ll focus on a single alphabet for clarity (think: the 26 letters), though nothing essential changes with multiple alphabets.

- Let the alphabet size be N arranged once and for all on a **ring** in a fixed cyclic order. We label letters by indices $i \in \{0, 1, \dots, N - 1\}$.

- A login consists of L **rounds**. In round r , the verifier chooses an independent **rotation** (a cyclic shift) $\theta_r \in \{0, \dots, N-1\}$. Intuitively, θ_r is determined by fresh entropy and trusted time; for our analysis we model it as an independent uniform random variable.
- After applying the rotation θ_r , the ring is partitioned into C **zones** of (as equal as possible) size about N/C . Think of painting the ring into C contiguous colored slices.
- The prover holds a **secret sequence** of letters $S = (i_1, \dots, i_L)$ (for example, “NOODLE”).
- The prover also holds a **private map of witnesses**, a *fixed* secret bijection

$$\hat{M} : \{0, 1, \dots, C-1\} \longrightarrow W,$$

from zone indices to a public set of **witness tokens** W (e.g., {UP, DOWN, LEFT, ...}). Observers can hear or see tokens in W , but do **not** know which token corresponds to which zone.

- Define the **zone index** of letter i under rotation θ by the function

$$q(i, \theta) \in \{0, \dots, C-1\},$$

which returns the contiguous slice (zone) containing $(i + \theta) \bmod N$. A concrete formula, assuming equal-sized zones when $C \mid N$, is

$$q(i, \theta) = \left\lfloor \frac{(i + \theta) \bmod N}{N/C} \right\rfloor.$$

(If $C \nmid N$, zones differ by at most one element; we handle that small detail later.)

What does a passive attacker observe? In every round r they can record the *visible layout* (which letters landed in which colored zone *that* round) and the *emitted witness token* $w_r = \hat{M}(q(i_r, \theta_r))$. They also see the final pass/fail. They do **not** know $\hat{M}$, and they cannot predict θ_r for future rounds.

The folklore claim we will dismantle is:

Folk Claim. “If an attacker records enough sessions (the layouts and the witnesses), they will eventually recover the secret sequence S and the private map $\hat{M}$.”

We will prove that the claim is **mathematically false** under the axioms of the protocol, and then we will explain why even desperate brute-force strategies are crushed by **physical limits** (capacity, energy, time).

Four axioms (the backbone)

To keep the argument clean, we isolate the assumptions that ENI6MA depends on. All are operational design choices that can be implemented and audited.

Axiom 1 (Independent, uniform rotations). For each round r , θ_r is independent of θ_s for $s \neq r$ and uniformly distributed over $\{0, \dots, N-1\}$.

Axiom 2 (Contiguous, near-equal zones). Each rotation θ_r induces a partition of the ring into C contiguous zones of sizes either $\lfloor N/C \rfloor$ or $\lceil N/C \rceil$.

Axiom 3 (Private bijection of witnesses). The witness map M is a fixed, per-prover bijection from zone indices $\{0, \dots, C-1\}$ to the public witness set W . It never leaves the prover's control and is unknown to observers.

Axiom 4 (Layout independent of the secret). The random choice of θ_r is statistically independent of the secret letters i_r . (Operationally: the verifier's randomness and time source does not depend on which letter the prover will check next.)

These are mild and realistic; they are exactly what the interface does.

Three lemmas (the engine room)

We build three lemmas that, chained together, destroy the folk claim.

Lemma 1 (Uniformity Through Ring Rotation).

Fix any letter i . If θ is uniform on $\{0, \dots, N-1\}$, then the random zone index $q(i, \theta)$ is uniform on $\{0, \dots, C-1\}$ when $C \mid N$; when $C \nmid N$, it is “near-uniform,” and the per-zone probability deviates from $1/C$ by at most $1/N$.

Idea of proof. The map $\theta \mapsto (i + \theta) \bmod N$ is a bijection on $\{0, \dots, N-1\}$, so $(i + \theta) \bmod N$ is uniform on ring positions. Partition the ring into C consecutive blocks (zones). If each block has size N/C (exact divisibility), the chance to land in any particular block is exactly $1/C$. If not, r blocks have size $\lfloor N/C \rfloor + 1$ and $C - r$ have $\lfloor N/C \rfloor$, so the gap between the largest and smallest zone probability is $\leq 1/N$.

Plain reading. A rotation is a rigid, global shift. As θ varies, *every* letter visits *every* zone equally often (exactly or within a tiny slack). A single “tick” of the dial moves **all** letters, so there's no way for one letter's zone frequency to stand out.

Lemma 2 (Independence across rounds and sessions).

If $\{\theta_r\}$ are independent (Axiom 1), then for any fixed secret letter i , the sequence $q(i, \theta_1), q(i, \theta_2), \dots$ is i.i.d. (independent and identically distributed) with the

uniform (or near-uniform) law from Lemma 1. Consequently, over multiple sessions, the entire multiset of zone indices produced by a fixed i remains *i.i.d.* with that same law.

Plain reading. Each round scrambles the ring afresh. There’s no “momentum” from one round to the next, so frequency and correlation attacks cannot accumulate signal.

Lemma 3 (Masking by the private bijection).

Let Z be any random variable on $\{0, \dots, C-1\}$. If $\hat{M}$ is a bijection $\{0, \dots, C-1\} \rightarrow W$ unknown to the observer and $W^* = \hat{M}(Z)$, then W^* and Z carry the same entropy but different labels; to an observer without $\hat{M}$, the distribution of W^* is indistinguishable from any permutation of Z ’s distribution.

In particular, if Z is uniform on $\{0, \dots, C-1\}$, then W^* is uniform on W . An observer who hears W^* cannot tell which zone index produced it.

Plain reading. The private witness map $\hat{M}$ is a perfect “label scrambler.” It turns a uniform (or near-uniform) signal over zone indices into the same uniform (or near-uniform) signal over tokens, but with labels that outsiders can’t interpret.

The main theorem (no information to learn)

Armed with those lemmas, we now **formalize** the headline claim.

Theorem (Zero mutual information for passive observation).

Fix the secret sequence $S = (i_1, \dots, i_L)$. Under Axioms 1–4, the entire passive transcript

$$T = ((\text{layout}_1, w_1), \dots, (\text{layout}_L, w_L))$$

has **zero mutual information** with S : $I(S; T) = 0$ when $C \mid N$ and $I(S; T) \leq O(L/N)$ in the non-divisible case (a tiny, removable slack).

Equivalently: *even infinitely many transcripts do not make S more predictable to an observer.*

Sketch of proof. By Axiom 4, layouts are generated independently of S (they depend only on θ_r). Condition on the realized layout in round r . For any fixed letter i_r , the zone index $q(i_r, \theta_r)$ is a **deterministic** function of the layout (it’s simply “which zone contains that letter, given what we see on the screen”). However, by Axiom 1 and Lemma 1, **marginalizing over the randomness that produced the layout**, the distribution of $q(i_r, \theta_r)$ is uniform (or near-uniform) and does **not** depend on i_r . Applying Lemma 3, the observed token $w_r = \hat{M}(q(i_r, \theta_r))$ is uniform (or near-uniform) over W , again **independent of** i_r . By independence across rounds (Lemma 2), the joint law of $(w_1, \dots, w_L)$

is i.i.d. uniform (or near-uniform) for any S . Since the layouts themselves are independent of S , the pair (layouts, tokens) has a distribution that **does not vary with S** . Therefore $I(S; T) = 0$ (or $\leq O(L/N)$ which vanishes with light padding to make $C \mid N$).

Translation. Even if you watch forever, the statistics of what you can see are the **same for every possible secret**. There is literally **no signal** in the transcript that points to the secret. No amount of data turns “no signal” into “some signal.”

This already refutes the folk claim in a strict information-theoretic sense. But we can also attack the claim from two more angles: **identifiability** and **sample complexity**.

A symmetry (identifiability) obstruction

There is a subtle, beautiful group-theoretic effect at play: **gauge freedom** in the labeling of zones and tokens.

- The private map $\hat{M}$ lives in the symmetric group S_C (all $C!$ permutations of C items).
- For any hypothesized secret S and any transcript of tokens $(w_1, \dots, w_L)$, there exists *some* bijection $\hat{M}$ that makes those tokens match the zones that would be correct for S **in hindsight**.
- To an outside observer, the pair $(S, \hat{M})$ is identifiable only up to this joint relabeling; **the same token transcript can be explained by many different $(S, \hat{M})$ pairs**.

This is an **identifiability obstruction**: even with unlimited transcripts, an attacker cannot settle on a *unique* pair $(S, \hat{M})$. The best they can do is produce a cloud of equally plausible explanations, all related by permutations. In statistics language, **the model is unidentifiable** from passive data.

This obstruction is not fragile, it’s structural. As long as the observer lacks a “ground truth” anchor that ties a specific token to a specific zone, the symmetric group acts transitively on the explanations. The only way to break that symmetry is to obtain **side information** about the private map $\hat{M}$ (which the protocol deliberately never exports).

Why classic attacks collapse

With the theorem and the symmetry in hand, the usual bag of tricks has nothing to grip:

1. Frequency analysis.

For a fixed letter i , the zone index $q(i, \theta)$ is uniform over $\{0, \dots, C - 1\}$.

After masking by $\hat{M}$, the token $w = \hat{M}(q)$ is uniform over W . Therefore, the frequency of tokens reveals *no* letter preferences. There is no “E appears a lot” signature to exploit.

2. **Correlation (bigram/trigram) attacks.**

The pair $(q(i_r, \theta_r), q(i_{r+1}, \theta_{r+1}))$ is a product of independent uniforms, so bigrams over tokens are also independent uniforms after masking. There is no “TH often follows T” footprint because each round’s geometry is independent.

3. **Layout-token cross-correlation.**

One might attempt to correlate the observed layout with the observed token: “When ‘A’ is in the red zone, does the user say UP more often?” But the token label “UP” is unrelated to “red” without knowing $\hat{M}$; and *even if* you guess a mapping for one session, the next round’s rotation moves **every** letter in lockstep, wiping your guess out statistically. (In exact terms: w is independent of i *even when conditioning on the layout*, the critical conditional-independence step inside the main theorem.)

4. **Repeated letters (e.g., the ‘OO’ in NOODLE).**

Two appearances of the same letter occur under independent rotations; their zones are independent draws, so they do not exhibit a detectable “repetition pattern.” The only structure is the verifier’s membership check, which is not visible to the observer beyond the final pass/fail.

This is the core of the design: **every round is a fresh coin flip** as far as the attacker is concerned.

Sample complexity: even the “slack” is unreachable

What about the small technicality when C does not divide N ? Then some zones are longer by one element. Does that whisper any information?

Let p_z be the probability that a fixed letter i lands in zone z . From Lemma 1,

$$\left| p_z - \frac{1}{C} \right| \leq \frac{1}{N}, \quad \sum_z p_z = 1.$$

A classical result in statistics says that to **distinguish** a perfectly uniform distribution from one with bias at most ε in total variation distance with confidence $1 - \delta$, you need on the order of

$$\Omega\left(\frac{1}{\varepsilon^2} \log \frac{1}{\delta}\right)$$

independent samples. Here $\varepsilon \lesssim 1/N$, so the sample complexity is $\Omega(N^2 \log(1/\delta))$.

- For $N = 26$ and $\delta = 10^{-6}$, this is on the order of tens of thousands of *independent* observations **per letter**, which you do not get in an authentication transcript.
- But in our setting, the observer never even sees the **zone index**, they see the **token** $w = \hat{M}(q)$. Since $\hat{M}$ is unknown, the bias is a bias on unlabeled categories; in effect, the labels can be permuted arbitrarily. That makes the statistical test even weaker: you do not know *which* category to test for a small surplus.

In short, even the vanishingly small slack is not practically learnable in this interaction model. And if one is still nervous, the implementation can trivially pad the alphabet (add dummies) so that N is a multiple of C , eliminating the slack *exactly*.

High-dimensionality: why “watch more” does not help

It’s easy to mistake **volume of data** for **information**. What you record each session is high-dimensional: a whole colored layout (many pixels) and a witness token. Surely patterns must lurk in such a big space?

Here is the trick: **the layout carries entropy that is independent of the secret** (Axiom 4). The per-session “movie” is drawn from a large distribution over ring rotations and zone partitions, but that distribution is **the same regardless of the secret**. In information-theory terms, the layouts are **pure noise** with respect to the secret variable S . The only “signal” that touches the secret is the membership event “the secret letter lies in the named zone,” but the only part of that event you’re allowed to see is the token w , which (by Lemmas 1–3) is i.i.d. uniform and thus also **pure noise** with respect to S .

The result is a **high-dimensional noise shell** around the secret: the more sessions you collect, the more independent noise you accumulate, without ever increasing the mutual information with S . This is not a hand-wavy metaphor; it is what $I(S; T) = 0$ means.

Practical limits (capacity, energy, time): the sledgehammer

Suppose someone ignores the theory and tries a blunt approach:

“I’ll just brute-force *everything*: enumerate secrets S , enumerate witness maps $\hat{M}$, check which pairs remain compatible with all my recordings, and pick the survivors.”

Even abstracting away the fact that **every** pair $(S, \hat{M})$ has the same transcript likelihood (so you end with a massive set of survivors anyway), this program falls off a cliff when you look at sizes.

- If the secret has length L over an alphabet of size N , then there are N^L candidate secrets.
- There are $C!$ possible private maps $\hat{M}$.
- So the naive hypothesis space is $N^L \cdot C!$. For typical human-usable parameters, say $N = 26$, $L = 8$, $C = 6$, that's $26^8 \cdot 720 \approx 2.08 \times 10^{12}$ hypotheses (trillions). And the search delivers no unique “winner” because of the identifiability obstruction: many $(S, \hat{M})$ pairs remain equally plausible.

But modern attackers have GPUs! A trillion is no longer unthinkable! So let us escalate to the protocol’s second line of defense: **cryptographic hardness** bound by a large entropy parameter λ (e.g., 256 or 512 bits) in the layout generation (the seed that drives the rotations). Even if the adversary tries to *predict* or *invert* the layouts to fake a future run, they face guessing a λ -bit value. The probability of success per try is $2^{-\lambda}$; for $\lambda = 256$, that’s about 10^{-77} .

At this point, physics steps in:

- A state-of-the-art **exascale** machine executes about 10^{18} operations per second. Run it for 30 years ($\approx 10^9$ seconds), that’s 10^{27} operations.
- Searching a $\lambda = 256$ space requires on average $2^{255} \approx 5.8 \times 10^{76}$ trials, **49 orders of magnitude** more than you have.
- Even if you had 10^9 such machines (a billion exascale computers) running for 30 years, you would still be short by **40+ orders of magnitude**.

And then there’s **energy**. Landauer’s principle says that erasing one bit dissipates at least $kT \ln 2$ joules of heat (where k is Boltzmann’s constant and T is temperature). Every realistic computer irreversibly erases many bits per operation, so combing through 2^{256} candidates would demand energy **far beyond** the total energy humanity can produce in a geological epoch. Whatever the precise constant factors, the conclusion is invariant: **your star dies before your search ends**.

In short: even if the math left a crack (it does not), the universe would shut the door.

The “replay” angle: why yesterday’s success is inert

A cousin of the folk claim is: “If I record a successful session, I can just replay it.” Not here.

Let $Z_r = q(i_r, \theta_r)$ be the zone index in round r for yesterday’s session, and let $w_r = \hat{M}(Z_r)$ be the emitted token. Today’s session uses fresh θ'_r ; the correct zone index for the same letter is $Z'_r = q(i_r, \theta'_r)$.

Key fact: If $\theta'_r \neq \theta_r$, then with overwhelming probability $Z'_r \neq Z_r$. So replaying yesterday's w_r simply points to the wrong zone today. Over L rounds, the chance that **all** rotations repeat perfectly (so a replay would pass) is $(1/N)^L$ (or, if the layout generation is keyed by a large λ , bounded by $2^{-\lambda}$), numbers so small they vanish. This is “anti-replay by construction”: changing the time-indexed rotations changes the manifold, which invalidates old transcripts.

This is a crucial philosophical shift. Traditional systems glue replay protection on the outside (timestamps, rate limiting). Here, replay is impossible **inside** the math.

A plain, “doorkeeper” intuition (without equations)

If the above felt heavy, here's a concrete picture:

- Imagine a circular board with the alphabet on it. Each round, the doorkeeper spins the board by a secret amount and then paints C slices (zones) on it. You know a word. To prove it, you **name a color** (zone) that currently contains your next letter.
- The eavesdropper hears what you say (a word like “UP”), and they can *see* how the board looked **that** round. But they **don't** know which color the word “UP” refers to, because that mapping is private to you and the doorkeeper.
- Even worse for the eavesdropper, next round the board is spun again. So any guess they had about “UP $\leftrightarrow$ red” is immediately scrambled. In fact, across many rounds, your spoken words are indistinguishable from random choices among the C zone names.
- Tomorrow, the board is spun differently again. Even if they play back your recorded words, they will be pointing at *yesterday's* zones, which no longer match *today's* positions of your letters.

The doorkeeper and performer always coordinate in the present tense. The past is just a souvenir.

Addressing Naive Objection / Challenges

Objection 1: “Okay, but I’m not passive. I’ll set up a fake verifier and trick the user into logging in through me. Then I control the layouts.”

That is a *man-in-the-middle* (MITM) attack. The protocol's default defense is that the *verifier* must check per-round acceptance against their own layouts. A fake verifier can't *complete* the proof with the real verifier without synchronizing layouts (which are keyed by entropy and time on the real side). There are also bilateral variants (“compiled twins”) where the user demands a witness

from the verifier; a MITM then has to fake *two* parties’ private maps coherently, far harder in practice.

Objection 2: “What if I try to reverse-engineer the private map $\hat{M}$ from many sessions?”

You cannot. The witness tokens are uniformly distributed across $\{1, \dots, C\}$ (after relabeling by $\hat{M}$) and independent across rounds. There is no statistical footprint to align a token to a specific zone. Any alignment you propose can be defeated by a different session’s rotations. Formally, the space of explanations is closed under the symmetric group on C labels. Unidentifiable.

Objection 3: “I’ll combine token frequencies with layout images and use machine learning to detect subtle patterns humans missed.”

All roads still lead to the mutual-information theorem. A perfect classifier cannot beat an information bound of zero. More data does not create information ex nihilo. Without side channels (timing, audio leakage about the private map, etc.), the input features are statistically independent of S . In ML terms: your model will learn the *layout generator*, not the user’s secret.

Objection 4: “But if I really watch a lot, like millions of sessions, surely the law of large numbers will reveal something.”

It will reveal that tokens are i.i.d. uniform, and that the layout generator behaves as designed. It will not reveal S . That is exactly what the law of large numbers says when the expectation is the same under every hypothesis.

The forging probability (and why it’s tiny on purpose)

The only strategy left is to *guess* the right zone each round. If you have no knowledge of the secret letter or the private map, your success per round is $1/C$. Over L rounds:

$$\Pr[\text{pass}] \leq C^{-L} + 2^{-\lambda}.$$

For human-friendly numbers like $C = 6$, $L = 6$, the geometric term is $6^{-6} \approx 1.6 \times 10^{-5}$. If the layout engine is keyed with $\lambda = 256$, the cryptographic term is effectively zero. With basic throttling (a handful of attempts per hour), *guessing* becomes a non-starter.

This bound is the protocol’s **soundness** guarantee: cheaters almost never pass.

Implementation hygiene (keeping the axioms true)

Everything above rests on Axioms 1–4. They are straightforward to uphold:

- Use a trusted time source and a cryptographically strong PRP/PRF to derive θ_r from a per-session seed; this gives independence and uniformity.
- Choose C to divide N (or pad N) so zones are equal sized.
- Keep the witness bijection $\hat{M}$ private and never reflect it in UI elements an observer could map (e.g., avoid labeling tokens with colors that match on-screen colors).
- Don't leak timing or micro-structure that correlates with the private map (constant-time UI responses, optional jitter).

These are all testable engineering practices, not leaps of faith.

A final, big-picture comparison

Traditional secrets are **objects** you can copy. This protocol turns a secret into a **capability** you must exercise *in the moment* against a freshly randomized world. That is why:

- **Recording** is harmless: you captured a performance in a world that no longer exists.
- **Statistics** are helpless: the distributions you observe are equal (or negligibly near-equal) for all secrets.
- **Brute force** is impossible: the parameter λ pushes the search beyond the computational and energetic means of our universe.
- **Replays** fail by design: changing the time re-tilts the world.

If you only remember one sentence, let it be this:

No number of passive observations can reveal a variable that the distribution does not depend on.

And in this scheme, the transcript distribution doesn't depend on the secret.

Epilogue: the axioms, lemmas, and theorem in one breath

- **Axiom 1–4:** Independent uniform rotations; contiguous near-equal zones; private bijection of witnesses; layouts independent of the secret.

- **Lemma 1 (Uniformization):** For any letter, the zone index after a random rotation is (near-)uniform over $\{0, \dots, C-1\}$.
- **Lemma 2 (Independence):** Across rounds and sessions, zone indices are i.i.d.
- **Lemma 3 (Masking):** Applying an unknown bijection to a uniform variable preserves uniformity and destroys label meaning.
- **Theorem (Zero mutual information):** The complete passive transcript, layouts plus tokens, has the same distribution for every secret S . So $I(S; T) = 0$ (or arbitrarily close to zero with trivial padding), and **no amount of observation helps**.

With that, the folk claim dissolves. Watching more does not unlock the secret, not mathematically, not statistically, not computationally, and not physically.

Setup (notation)

- Fix an alphabet of size N laid out once and for all in a **ring** (a cycle) in a *fixed* order. Label letters by indices $i \in \{0, 1, \dots, N-1\}$ around the ring.
- Each *round* r the verifier chooses a **rotation** (cyclic shift) $\theta_r \in \{0, 1, \dots, N-1\}$. Think of θ_r as derived from entropy+time and modeled as independent uniform random variables across rounds.
- After rotating by θ_r , the ring is **foliated** into C contiguous **zones** of (as equal as possible) size $\approx N/C$. Formally, define the zone map

$$q(i, \theta) = \left\lfloor \frac{(i + \theta) \bmod N}{\lfloor N/C \rfloor \text{ or } \lceil N/C \rceil} \right\rfloor \in \{0, 1, \dots, C-1\}.$$

(When N is divisible by C , each zone has size exactly N/C , and the denominator above is simply N/C .)

- The prover’s **private map of witnesses** is a secret bijection $\hat{M} : \{0, \dots, C-1\} \rightarrow W$ from zone indices to a public set of witness tokens W (e.g., words like “UP/LEFT/...”, icons, etc.). Observers hear the token $w = \hat{M}(q)$ but **do not** know $\hat{M}$.

The only things an attacker/observer can do are: record the visible ring layout (which letters fell in which zone *that round*) and record the emitted witness token w .

Key lemma: rotations make zone membership uniform

Lemma 1 (Uniformization by rotation).

Fix any letter index i . If θ is uniformly random on $\{0, \dots, N-1\}$, then $q(i, \theta)$ is **uniform** on $\{0, \dots, C-1\}$ whenever N is divisible by C ; when N is not divisible by C , the distribution is *near*-uniform and the per-zone bias is at most $1/N$.

Proof (sketch).

The map $\theta \mapsto (i + \theta) \bmod N$ is a bijection of $\{0, \dots, N-1\}$ to itself, so a uniform θ makes $(i + \theta) \bmod N$ uniform on $\{0, \dots, N-1\}$. Partition that set into C consecutive blocks (the zones). If N is divisible by C , each block has size exactly N/C , so the zone index is uniform with probability $1/C$ for each zone. If $N = qC + r$ with $0 < r < C$, then r zones have size $q+1$ and $C-r$ have size q , so the largest probability is $(q+1)/N$ and the smallest is q/N , giving max deviation $\leq 1/N$.

Interpretation. A rotation is the same “increment” applied to **every** letter index. That *uniformly shifts the entire ring*. As θ varies, each fixed letter falls into each zone equally often (exactly, or up to a tiny $1/N$ slack). This is what you were pointing at when you said: “a single change of the elements causes a uniform change across all other characters.”

Consequences for attacks

1) Frequency attacks fail

A classic frequency attack tries to learn which *plaintext* letter is which by watching which *symbol* appears most often. Here, what the observer sees each round is only the **zone** of the current secret letter (and even that is masked by $\hat{M}$; see below). For any fixed plaintext letter i , by Lemma 1,

$$\Pr[q(i, \theta) = z] = \frac{1}{C} \quad (\text{exact if } C \mid N; \text{ within } 1/N \text{ otherwise})$$

for each $z \in \{0, \dots, C-1\}$. Therefore, **zone frequencies are identically distributed for all letters**. The most common plaintext letter (“E”, say) does **not** make any zone (or any witness token) appear more often than it would for a rare letter. Observed frequencies carry no information about the underlying letter distribution.

If rounds use independent θ_r , then for a fixed letter i the sequence $q(i, \theta_1), q(i, \theta_2), \dots$ is i.i.d. uniform on $\{0, \dots, C-1\}$. So even very long observation windows do not give the attacker a statistical edge.

2) Correlation (bigram) attacks fail

Suppose an attacker tries to exploit correlations like “T is often followed by H”. In round r , the zone is $q(i_r, \theta_r)$ for the r -th letter i_r of the secret. Because θ_r and θ_{r+1} are independent, the pair $(q(i_r, \theta_r), q(i_{r+1}, \theta_{r+1}))$ has a **product** distribution: each component is (near-)uniform and independent of the other. Thus, the joint distribution over consecutive rounds is uniform on $\{0, \dots, C-1\}^2$ (up to the same $1/N$ slack when $C \nmid N$). No bigram signature survives.

Even if someone hypothesizes the same letter repeats (e.g., double “O”), the two appearances happen under independent rotations, so their zones are independent draws; there is no “repeated-letter” footprint to correlate.

Why a *ring* (fixed order + random rotation) is enough

You might ask whether we need a fresh *random permutation* of letters each round. Surprisingly, **no**, for these leakage questions, using the **cyclic subgroup** (rotations) already gives the same protection:

- The group of rotations $\{x \mapsto x + \theta \bmod N\}$ acts **transitively** on positions: any letter visits every position equally often as θ varies.
- Because zones are defined by contiguous blocks of positions, transitivity implies the exact **uniformity of zone membership** (again, exact when $C \mid N$, near-exact otherwise).
- Crucially, an increment $\theta \mapsto \theta + 1$ moves **every** letter forward by one position. That “rigid body” motion guarantees that **any local change is a global change**, you cannot bias one letter’s zone without moving all letters in lockstep. This “all-or-nothing” shift is what defeats attempts to track or pin a particular letter by watching small drifts.

A fully random permutation is *overkill* here: it gives the same marginal law for zone membership but is harder for humans to scan. The fixed-order ring plus random rotation preserves human legibility while already achieving the probability-theoretic uniformity that kills frequency/correlation attacks.

Private witness map $\hat{M}$ finishes the job

Observers do not see the numeric zone z ; they hear a **witness token** $w = \hat{M}(z)$, where $\hat{M}$ is a **private bijection** known only to the prover/verifier.

- If Z is uniform on $\{0, \dots, C-1\}$ and $W = \hat{M}(Z)$, then W is uniform on W (basic property of bijections).

- Because $\hat{M}$ is unknown to observers, they cannot align “UP/LEFT/...” to specific colors/regions on the display. Even if they record many rounds, they only accumulate uniform samples of W , indistinguishable from what *any* secret would have produced.

Formally, conditioned on any plaintext letter i , W has the same distribution. Thus the **mutual information** $I(i; W)$ is 0 (exactly when $C \mid N$; within negligible $\tilde{O}(1/N)$ otherwise), and likewise for tuples across rounds because of independence of the θ_r .

Edge case: when N is not a multiple of C

As noted, when $N = qC + r$ with $0 < r < C$, some zones have one extra symbol. Two points matter:

1. The maximum bias per zone is at most $1/N$. For typical N (e.g., $N = 26$) and small C (say $C \leq 8$), this is tiny.
2. Rotations are independent each round and transcripts are short. Estimating a $1/N$ -level bias would need far more samples than any authentication transcript provides, and the unknown $\hat{M}$ still masks which physical zone a token corresponds to.

If desired, implementations can pad alphabets (e.g., add dummies to make N divisible by C) so Lemma 1 holds **exactly**.

Putting it together

- **Uniformity:** For any letter, its zone after a random rotation is uniform on $\{0, \dots, C - 1\}$.
- **Independence:** Across rounds, zones are independent because rotations are independent.
- **Ring suffices:** Rotations (the cyclic subgroup) already produce the needed uniform/independent laws.
- **Masking:** The private bijection $\hat{M}$ turns uniform zones into uniform witness tokens, hiding which zone was meant.
- **No signal to mine:** Frequencies, bigrams, and cross-round correlations of what observers can record are identical for *all* secrets.

Hence, **recording the ring layouts and the witnesses gives an attacker no statistical leverage**: every letter is equally likely to induce every witness, every round, and the private map $\hat{M}$ severs any link between what’s heard and what zone was actually named.

Conclusion: Why ENI6MA Stays Secure, Not Mathematically, Not Statistically, Not Computationally, and Not Physically

If this essay had to be distilled to a single line, it would be this: **ENI6MA changes the nature of a secret from a reusable object into a per-session performance played on a stage that is rebuilt every round.** From that single design choice, the four “nots” follow, mathematical, statistical, computational, and physical, and together they close off the ordinary paths an attacker would take. This concluding section pulls the threads tight, showing how the parts cohere, where the guarantees come from, what would be required to break them, and why those requirements are out of reach in any world we can actually build.

Not mathematically.

The mathematical claim is about *dependence*: does the distribution of what an attacker can observe depend on the victim’s secret? In ENI6MA, the answer is no. Each round begins with a fresh, independent rotation of a fixed-order ring and a foliation into zones. For any letter, the chance to land in any zone is equal (or negligibly close to equal when the alphabet doesn’t divide perfectly), and across rounds those placements are independent. The witness the user emits is a relabeling of that zone via a private bijection that never leaves the user. The result is that the observable transcript, layouts, tokens, pass/fail, has the *same* distribution for every possible secret. In information-theory language, the mutual information between the secret and the transcript is zero (or arbitrarily close to zero with trivial padding). When the output of a channel does not depend on its input, no theorem, no algorithm, and no devilishly clever trick can recover that input. The premise of inference is missing. Thus, **mathematics rules out recovery by observation.**

Not statistically.

Where mathematics states “there is no signal,” statistics asks, “could a signal emerge with enough samples?” Here, too, the answer is no. Frequency analysis fails because rotations spread every letter evenly across zones, flattening counts. Correlation attacks fail because each round’s geometry is fresh and independent, breaking bigram or higher-order structure. Cross-modal attacks that try to align tokens to visual zones fail because the private token→zone map never appears in the observable world, and any tentative alignment is scrambled by the next rotation. Even the tiny edge case, zones differ by at most one element when the alphabet doesn’t split cleanly, produces biases so small, and on unlabeled categories, that they are unmeasurable within any reasonable number of login rounds; and they can be eliminated entirely by padding the alphabet. Machine learning does not change the story: models cannot learn a dependency that is absent. With more data, you estimate the same flat distribution more precisely,

you do not conjure structure where none exists. Thus, **statistics cannot manufacture evidence out of independent, uniformly distributed noise.**

Not computationally.

If you cannot learn the secret by watching, perhaps you can **guess** it. ENI6MA is structured so that the guessing game is hostile from the first moment. The search space for secrets grows exponentially with length; the space of private maps is factorial in the number of zones. Their product explodes quickly, even for friendly human-scale parameters. Meanwhile, the per-attempt success probability is engineered to be tiny: if there are C zones and L rounds, blind success is about $(1/C)^L$. Rate limits (and lockouts) turn that tiny probability into a practical impossibility. And if an attacker hopes to get around geometry by predicting or inverting the layout engine, they run into cryptographic hardness: the rotations derive from a seed with hundreds of bits of entropy. No practical enumeration of secrets, maps, and seeds exists; no clever pruning makes the combinatorics benign; and no shortcut appears because the protocol exports no hooks to grab onto. Even wildly optimistic hardware projections fall orders of magnitude short of the work needed. In short, **computation cannot bridge the chasm between the available resources and the required search.**

Not physically.

Computation rides on energy and time. Even if one ignores engineering overhead and imagines perfect parallelization, exploring cryptographic-scale spaces demands numbers of operations so vast that the *minimum* heat they would dump (by fundamental thermodynamic limits) becomes planetary, and the *minimum* time to perform them surpasses plausible lifetimes of data centers, organizations, or individuals. Adding more machines only brings forward the energy wall and heat dissipation crisis. Quantum speedups (e.g., Grover’s algorithm) do not rescue the attacker; a quadratic improvement against an exponential target still leaves an exponential that is far beyond reach when parameters are chosen sensibly. Thus, **physics itself vetoes the fantasy that raw force can make up for absent information.**

The four locks reinforce one another.

These are not four independent lines of defense that might fail one at a time; they are one structure seen from four angles. The math of independence and uniformity generates the statistical flatness; that flatness leaves nothing for computation to grind; and the gulf between required and available computation is widened by physical law. Pull on any thread and you feel the whole fabric.

Replay and “record-to-reuse” are dead on arrival.

Because each session’s geometry is keyed to new randomness and time, yesterday’s valid path is incompatible with today’s stage. A pristine recording of a successful proof does not reassemble the secret; it does not even produce a sequence that would be accepted again. It is a fossil, evidence that a performance happened once, inert thereafter. This is a categorical difference from passwords and static credentials, where a captured secret remains a secret.

What about active meddling?

Man-in-the-middle games, relays, and phishing pages all depend on synchronizing two worlds while lacking the private map and the verifier’s time-indexed geometry. They cannot maintain round-by-round consistency. Variants with mutual witnesses make the imposter’s task strictly harder, forcing them to counterfeit *both* sides’ private legends coherently. In practice, the same ingredients, fresh randomness, hidden legend, verifier-side checking, break the symmetry that active attackers try to exploit.

Assumptions and hygiene, keeping the axioms true.

Every real system rests on choices. ENI6MA’s guarantees assume four that are simple and auditable: per-round rotations are fresh and independent; zones are contiguous and near-equal (or exactly equal with padding); the token→zone map remains private; and the layout engine is independent of the secret. Implementation hygiene (constant-time UI, no accidental leakage of the legend via colors or CSS or audio cues, rate limiting, solid randomness and trusted time) preserves those axioms. These are not fragile, esoteric requirements; they are ordinary engineering practices aligned with the protocol’s intent.

Security without custodial risk.

A powerful but sometimes overlooked consequence is *who holds what*. The verifier keeps receipts and non-secret parameters, not reusable credentials; the user keeps the private legend. Auditors can verify that each round’s witness really matched the laid-out zone for the committed letter without learning the secret or the legend. That yields transparency without creating treasure troves for attackers.

The right mental model to carry forward.

If you think of authentication as possession, “show me the string I can copy”, ENI6MA will feel alien. The better model is performance: *can you stay in step with tonight’s rhythm using your private choreography, while the floor keeps rotating?* A spectator can watch, cheer, and film, but the film won’t help them dance tomorrow. A stubborn onlooker can try to improvise the routine, but the tempo and arrangement will change again, and the door staff will give them only a few shots before the band moves on. The shift from possession to performance is what makes “watching more” and “guessing more” both sterile strategies.

Closing claim.

So when we say ENI6MA “cannot be hacked,” we are not invoking bravado. We are describing a protocol whose *observable outputs do not depend on the secret*; whose *statistics stay flat no matter how long you watch*; whose *guessing game is deliberately stacked so that luck is both mathematically rare and operationally throttled*; and whose *computational requirements to brute-force the unseen internals outstrip the energy and time available in any realistic universe*. Under those conditions, and with the straightforward operational hygiene that keeps them

in place, the everyday routes to compromise are closed. Not mathematically, not statistically, not computationally, and not physically.

Appendix: Equations

E1. Zone index (general form).

$$q(i, \theta) = \left\lfloor \frac{(i+\theta) \bmod N}{\lfloor N/C \rfloor \text{ or } \lceil N/C \rceil} \right\rfloor \in \{0, \dots, C-1\}.$$

Maps letter i on the fixed ring, after rotation θ , to its zone (slice) index.

E2. Zone index (equal-sized slices case).

$$q(i, \theta) = \left\lfloor \frac{(i+\theta) \bmod N}{N/C} \right\rfloor.$$

Same as E1 when $C \mid N$; all slices have size N/C .

E3. Private witness map (bijection).

$$\hat{M} : \{0, \dots, C-1\} \rightarrow W.$$

A secret, per-user relabeling from zone indices to public witness tokens (e.g., “UP/LEFT/...”).

E4. Observed token each round.

$$w_r = \hat{M}(q(i_r, \theta_r)).$$

What outsiders hear/see each round: the user’s token for the zone containing the current secret symbol.

E5. Per-letter zone uniformity (frequency flatness).

$$\Pr[q(i, \theta) = z] = \frac{1}{C} \text{ (exact if } C \mid N, \text{ within } 1/N \text{ otherwise).}$$

Every letter lands in every zone equally often across random rotations (near-exact if not divisible).

E6. Transcript object.

$$T = ((\text{layout}_1, w_1), \dots, (\text{layout}_L, w_L)).$$

Everything a passive observer can record over L rounds: each round’s layout plus the emitted token.

E7. Zero-information claim.

$$I(S; T) = 0 \text{ (exact if } C \mid N; \text{ otherwise } I(S; T) \leq O(L/N)).$$

The transcript distribution does not depend on the secret S ; observation yields no information.

E8. Independence across rounds (iid statement).

$\{q(i, \theta_1), q(i, \theta_2), \dots\}$ is i.i.d. uniform on $\{0, \dots, C-1\}$ when θ_r are independent.

Fresh rotations make zone membership independent from round to round.

E9. Forgery probability (soundness bound).

$$\Pr[\text{pass}] \leq C^{-L} + 2^{-\lambda}.$$

Blind-guess success across L rounds with C zones plus negligible chance to guess the entropy/seed.

E10. (From the formalism) Acceptance and membership.

$$M(s_r, W_{z_r}^r) = 1_{\{s_r \in W_{z_r}^r\}}, \quad \Lambda = \bigwedge_{r=1}^L M(s_r, W_{z_r}^r).$$

Each round checks membership of the current secret symbol in the named zone; accept iff all rounds pass.

E11. (From the formalism) Entropy–time–offset pipeline.

$\kappa_{r,j} = ((\tau \bmod U) + rM + j) \bmod U$, $\Theta_{\alpha_j}^{(r)} = \text{block}_{r,j} \bmod |A_j|$.
Deterministically mixes high-entropy seed E and time τ into per-round ring rotation offsets.

E12. (From the formalism) Zone witness construction.

$$W_z^r = \circ_{a=1}^M \alpha_{a,z}^r, \quad M(s_r, X_r) = \begin{cases} 1 & s_r \in X_r \\ 0 & \text{else} \end{cases}.$$

Concatenates the z -th slices across rings to define the zone set used in the membership test.

E13. (From the formalism) PoK theorem summary, including soundness.

$\Pr[\text{forge}] \leq C^{-L} + 2^{-\lambda}$ with completeness and passive ZK.

Formal statement of completeness/soundness/zero-knowledge for ENI6MA’s interaction.

Appendix: Symbol map

- N — Alphabet size (letters arranged on a fixed-order ring).
- C — Number of zones (slices) the ring is partitioned into each round.
- i — Index of a particular letter on the ring.
- θ, θ_r — Rotation (cyclic shift) applied to the ring; θ_r is the independent per-round rotation.
- $q(i, \theta)$ — Zone index in $\{0, \dots, C-1\}$ of letter i after rotation θ .
- z, z_r — A specific zone index; z_r is the user’s declared zone in round r .
- W — The public set of witness tokens (labels like “UP”, “LEFT”, etc.).
- $\hat{M}$ — Private bijection mapping zone indices to witness tokens.
- w_r — The token emitted in round r : $w_r = \hat{M}(q(i_r, \theta_r))$.
- L — Number of rounds in a login ceremony / proof.
- $S = (i_1, \dots, i_L)$ — The secret sequence of letter indices (user’s commitment).
- T — Full passive transcript $((\text{layout}_r, w_r))_{r=1}^L$.
- $I(S; T)$ — Mutual information between the secret and the transcript.
- $M(\cdot, \cdot)$ — Membership predicate: 1 if a symbol lies in the named zone’s witness set, else 0.
- Λ — Conjunctive acceptance accumulator over rounds; accept iff $\Lambda = 1$.

- λ — Entropy/security parameter (bits) protecting layout generation (e.g., 256–512).
 - $\text{Pr}[\text{pass}]$ — Probability a forger is accepted in a blind attempt.
 - E — High-entropy base integer (seed) used to derive rotations.
 - τ — Time coefficient (high-resolution clock) mixed with entropy per session.
 - $\kappa_{r,j}$ — Rosario Modulo Index selecting which entropy slice drives round r , alphabet j .
 - $\Theta_{\alpha_j}^{(r)}$ — Rotation offset for alphabet j in round r .
 - A_j — j -th alphabet ring; $|A_j|$ is its size.
 - $\alpha_{a,z}^r$ — The z -th slice of alphabet a after rotation in round r .
 - W_z^r — Zone witness set in round r (concatenated slices across alphabets).
-

Appendix: Axioms, lemmas, and proofs — one-sentence summaries

Axiom 1 (Independent, uniform rotations).

Each round’s ring rotation θ_r is an independent, uniform cyclic shift, ensuring no cross-round statistical footprint to mine.

Axiom 2 (Contiguous, near-equal zones).

After rotating, the ring is partitioned into C consecutive zones of (near-)equal size so that zone membership is balanced.

Axiom 3 (Private bijection of witnesses).

The user’s secret map $\hat{M}$ bijects zone indices to public tokens and is never observable, so labels seen by outsiders are unanchored.

Axiom 4 (Layout independent of the secret).

The random layout (rotations/foliation) is generated independently of the user’s secret letters, preventing leakage through correlations.

Lemma 1 (Uniformization by rotation).

For any fixed letter i , random rotation makes its zone $q(i, \theta)$ uniform over $\{0, \dots, C-1\}$ (exact if $C \mid N$, within $1/N$ otherwise).

Lemma 2 (Independence across rounds).

With independent rotations, the sequence of zones a fixed letter occupies across rounds is i.i.d., destroying frequency/correlation signals.

Lemma 3 (Masking by the private bijection).

Applying an unknown bijection $\hat{M}$ to zone indices preserves uniformity and strips label meaning, so observed tokens remain indistinguishable.

Main Theorem (Zero mutual information for passive observation).

The complete passive transcript T (layouts + tokens) has the same distribution for every secret S , so $I(S; T) = 0$ (or $O(L/N)$ if $C \nmid N$), i.e., watching more never helps.

Soundness Bound (Forgery probability).

A blind attacker's chance to pass is at most $C^{-L} + 2^{-\lambda}$, combining per-round zone guessing with negligible entropy guessing.

(From the formalism) Proof-of-Knowledge Theorem.

ENI6MA is complete, sound with $\Pr[\text{forge}] \leq C^{-L} + 2^{-\lambda}$, and passive zero-knowledge (transcripts simulatable without the secret/morphism).

(From the formalism) Anti-Replay Theorem.

Changing the time index τ necessarily changes at least one round's rotation offset, so old transcripts cannot validate in new sessions.

Bibliography

1. Shannon, C. E. (1949). *Communication Theory of Secrecy Systems*. **Bell System Technical Journal**, 28(4), 656-715.
<https://doi.org/10.1002/j.1538-7305.1949.tb00928.x>
2. Cover, T. M., & Thomas, J. A. (2006). *Elements of Information Theory* (2nd ed.). **Wiley**.
<https://onlinelibrary.wiley.com/doi/book/10.1002/047174882X>
3. Maurer, U. (1994). *Information-Theoretic Cryptography*. In **CRYPTO'94** (LNCS 839, pp. 287-305). Springer.
https://doi.org/10.1007/3-540-48658-5_27
4. Goldreich, O. (2004). *Foundations of Cryptography, Vol. 2: Basic Applications*. **Cambridge University Press**.
<https://doi.org/10.1017/CBO9780511721656>
5. Lindell, Y. (2022). *Secure Multiparty Computation and Secret Sharing*. **Cambridge University Press**.
<https://doi.org/10.1017/9781108957295>
6. Katz, J., & Lindell, Y. (2020). *Introduction to Modern Cryptography* (3rd ed.). **CRC Press**.
<https://doi.org/10.1201/9780429492493>
7. Renner, R., & Wolf, S. (2004). *Smooth Rényi Entropy and Applications in Cryptography*. **IEEE ISIT**.
<https://doi.org/10.1109/ISIT.2004.1365483>
8. Bennett, C. H., Brassard, G., Crépeau, C., & Maurer, U. M. (1995). *Generalized Privacy Amplification*. **IEEE Trans. Inf. Theory**, 41(6),

- 1915-1923.
<https://doi.org/10.1109/18.476316>
9. Landauer, R. (1961). *Irreversibility and Heat Generation in the Computing Process*. **IBM J. Res. Dev.**, 5(3), 183-191.
<https://doi.org/10.1147/rd.53.0183>
 10. Bennett, C. H. (2003). *Notes on Landauer's Principle, Reversible Computation, and Maxwell's Demon*. **Stud. Hist. Philos. Mod. Phys.**, 34(3), 501-510.
[https://doi.org/10.1016/S1355-2198\(03\)00039-X](https://doi.org/10.1016/S1355-2198(03)00039-X)
 11. Bérut, A., et al. (2012). *Experimental Verification of Landauer's Principle*. **Nature**, 483, 187-189.
<https://doi.org/10.1038/nature10872>
 12. Grover, L. K. (1996). *A Fast Quantum Mechanical Algorithm for Database Search*. **STOC '96**, 212-219.
<https://doi.org/10.1145/237814.237866>
 13. Nielsen, M. A., & Chuang, I. L. (2010). *Quantum Computation and Quantum Information* (10th Anniv. ed.). **Cambridge University Press**.
<https://doi.org/10.1017/CBO9780511976667>
 14. Brassard, G., Høyer, P., Mosca, M., & Tapp, A. (2002). *Quantum Amplitude Amplification and Estimation*. **Contemporary Mathematics**, 305, 53-74.
<https://doi.org/10.1090/conm/305/05215>
 15. Aaronson, S. (2013). *Quantum Computing Since Democritus*. Cambridge University Press.
<https://doi.org/10.1017/CBO9780511977879>
 16. Mosca, M. (2018). *Cybersecurity in an Era with Quantum Computers*. **IEEE Security & Privacy**, 16(5), 38-41.
<https://doi.org/10.1109/MSEC.2018.053711661>
 17. Alagic, G., et al. (2020). *Status Report on the Second Round of the NIST Post-Quantum Cryptography Standardization Process*. NISTIR 8309.
<https://doi.org/10.6028/NIST.IR.8309>
 18. Bernstein, D. J., & Lange, T. (2017). *Post-Quantum Cryptography*. **Nature**, 549, 188-194.
<https://doi.org/10.1038/nature23461>
 19. Bellare, M., & Rogaway, P. (1993). *Random Oracles are Practical: A Paradigm for Designing Efficient Protocols*. **ACM CCS '93**, 62-73.
<https://doi.org/10.1145/168588.168596>

20. Damgård, I., & Pedersen, T. P. (1993). *Proofs of Ownership of Secrets. CRYPTO '93* (LNCS 773), 82-96.
https://doi.org/10.1007/3-540-48329-2_8
21. Feige, U., Fiat, A., & Shamir, A. (1988). *Zero-Knowledge Proofs of Identity. Journal of Cryptology*, 1(2), 77-94.
<https://doi.org/10.1007/BF02351717>
22. Blum, M., Feldman, P., & Micali, S. (1988). *Non-Interactive Zero-Knowledge and Its Applications. STOC '88*, 103-112.
<https://doi.org/10.1145/62212.62222>
23. Chaum, D., & van Antwerpen, H. (1990). *Undeniable Signatures. CRYPTO '89*, 212-216.
https://doi.org/10.1007/0-387-34805-0_19
24. Bellare, M., Kilian, J., & Rogaway, P. (2000). *The Security of the CBC MAC. J. Computer & System Sciences*, 61(3), 362-399.
<https://doi.org/10.1006/jcss.1999.1694>
25. Boneh, D., & Shoup, V. (2020). *A Graduate Course in Applied Cryptography*.
<https://crypto.stanford.edu/~dabo/cryptobook/>
26. Alwen, J., & Serbinenko, V. (2015). *High-Entropy Proofs and Information-Theoretic Security. TCC 2015*, 83-102.
https://doi.org/10.1007/978-3-662-46494-6_5
27. Arute, F., et al. (2019). *Quantum Supremacy Using a Programmable Superconducting Processor. Nature*, 574, 505-510.
<https://doi.org/10.1038/s41586-019-1666-5>
28. Preskill, J. (2018). *Quantum Computing in the NISQ Era and Beyond. Quantum*, 2, 79.
<https://doi.org/10.22331/q-2018-08-06-79>
29. Chatterjee, S., & Nisan, N. (2021). *The Physical Limits of Computation and Implications for Cryptography. Found. Trends Theor. Comp. Sci.*, 16(3), 173-260.
<https://doi.org/10.1561/04000000096>
30. Vazirani, U., & Vidick, T. (2019). *Fully Device-Independent Quantum Key Distribution. Communications of the ACM*, 62(4), 133.
<https://doi.org/10.1145/3318172>