

What If the Slop Is Mine?

An AI-assisted Voynich rabbit hole becomes a self-audit of verification debt, external markers, and when a research system must refuse to publish.

DENNIS HEDEGREEN OPEN V1.0 [HTTPS://HEDEGREENRESEARCH.COM/ARTICLES/WHAT-IF-THE-SLOP-IS-MINE/](https://hedegreenresearch.com/articles/what-if-the-slop-is-mine/)

I did not open ChatGPT because I thought I had solved the Voynich manuscript. I opened it because a video about people who thought they had solved the Voynich manuscript made me wonder whether Hedegreen Research was a more elaborate version of the same problem.

The video was [*AI killed the Voynich theory*](https://www.youtube.com/watch?v=V_2s0PwAQtw) (https://www.youtube.com/watch?v=V_2s0PwAQtw), published by Koen Gheuens on his Voynich Talk channel. Gheuens is a long-serving moderator at Voynich Ninja, an online forum for people studying the manuscript. He described a field that had changed. Amateur solutions were not new. For more than a century, people had arrived convinced that they had translated the unreadable book. What was new was the speed, similarity, packaging, and distribution of the theories.

A person could mention Voynich to a chatbot. The chatbot could help construct a theory, give it academic language, generate a paper, produce references, recommend a repository, suggest a DOI, and point the user toward the forum, libraries, researchers, and journals. The result did not remain inside one person's private speculation. It arrived as work for other people.

The theory might take hours to generate. Checking it could take days. That asymmetry is not new. It is often described through Brandolini's bullshit-asymmetry principle: producing a weak or false claim is far cheaper than correcting it. Generative AI industrialises the production side without automatically reducing the cost of verification.

Gheuens described invented references, invented experiments, invented strings supposedly taken from the manuscript, and papers that looked enough like research to demand attention from people who knew the field. His community had reached the point where AI-assisted "solutions" were being held outside normal discussion.

I understood the frustration. Then I recognised some of the furniture.

Hedegreen Research uses AI heavily. I work with frameworks, protocols, models, PDFs, repositories, data, tools, and outreach. I have deposited work with DOI records. I publish at a speed I could not maintain alone. I am an independent operator outside a conventional university or research institution. From a distance, some of the external signals are uncomfortably similar.

That does not make the work equivalent. But it also does not prove that it is different.

The doubt came first

The doubt did not appear after ChatGPT helped me build a Voynich idea. It was already there when I watched the video, which is why the video mattered.

I have built a large amount of work in a short period. Hedegreen Research now contains articles, tools, datasets, PDFs, internal systems, source structures, visual language, publication workflows, and an expanding databank of relations. I can show that the work exists. I can show when it was produced, how versions changed, and in many cases which sources were used.

That is proof of work. It is not automatically proof of correctness.

A DOI proves that an object has a stable identifier. It does not prove that the object is right. A public repository proves that a file can be found. It does not prove peer review. A large archive proves production. It does not prove judgement. A reply from a researcher proves contact. It does not prove agreement. These distinctions matter because they are exactly the sort of external signals that can make weak work look stronger than it is.

I attended an EU civic-tech hackathon in Brussels. That is part of my real history. But participation followed an application or registration process, and I do not know how closely anyone examined Hedegreen Research before I arrived. I can honestly say that I attended, built, spoke with people, and received useful signals. I cannot turn that into an institutional endorsement that I was never given.

Being open means being open about that distinction too. Transparency is not a certificate of quality. It is a condition that makes inspection possible.

The uncomfortable question remained:

How do I know whether AI is helping me perform real research work, or merely making my intuitions faster, larger, and more research-shaped?

So I downloaded the transcript. Then I told ChatGPT not to write an article. I only wanted to talk.

A human idea enters the machine

What follows is not offered as a contribution to Voynich research. It is the specimen used to examine what AI does to an intuition before the underlying work has been done.

I had seen one documentary about the Voynich manuscript years earlier. Beyond that, I knew little about the current field.

My first association was not cryptography. It was the narrower possibility of a private system of meanings: something structured from within and almost impossible to reconstruct from outside. That association was not evidence, and it created no responsible route to diagnosing an unknown medieval author.

Then the first resistance arrived. I learned that the medievalist and palaeographer Lisa Fagin Davis had identified five scribal hands in the manuscript. Whatever produced the object was not safely reducible to a story about one isolated mind. If there had been a guiding private system, it would still have needed a collaborative environment: a workshop, household, group of pupils, assistants, scribes, or another community capable of carrying it.

Then Tolkien appeared.

Imagine that Tolkien had written only for his children. Imagine that the published novels never existed. Imagine that five hundred years later, researchers found maps, genealogies, invented scripts, fragments in Quenya and Sindarin, contradictory versions of stories, drawings, symbols, and notes referring to events nobody else remembered.

They might prove that the languages had structure. They still might not know what kind of object they had found. A lost religion? A record of an unknown people? A code? A private cosmology? A work of fiction without its missing entrance? The missing person, household, or circle could have been part of the key.

That led to a different Voynich question. Not: *What does each sign translate to?* But: *What kind of person, group, and environment could have produced an object that remained regular after its meaning became inaccessible?*

Even a highly original private world has local inputs. Tolkien could build Middle-earth, but he could not remove the languages, landscapes, stories, and histories that had entered him. The more original the system, the more interesting its accidental leaks may become.

If the Voynich author invented plants, symbols, categories, and language, the invention could still contain traces of nearby manuscripts, buildings, containers, medical practices, drawing habits, local plants, religious conflicts, trade routes, and ways of organising knowledge.

You may not be able to translate the private world. You may still be able to locate some of its ordinary inputs. At this point, I had an intuition. Then AI turned it into a programme.

The result was useful. That was the problem.

Within a few hours, the conversation expanded across fields I do not command: dating and provenance, palaeography, parchment production, pigments, manuscript workshops, related texts, historical maps, local practices, and possible material-analysis methods.

Some of those directions may be useful. Some may be weak. Some may already have been investigated. Some may be technically possible but historically meaningless. Some may have been overstated by the AI in the conversation.

That distinction did not slow the expansion. The idea had acquired departments.

A human researcher could have done much of this work. A competent research assistant could search the literature, identify relevant disciplines, locate specialist methods, collect candidate regions, and build a first map of the problem. A better specialist would do it more carefully, reject false connections earlier, and understand the field's disputes in ways a language model does not.

But the AI did perform labour. It found vocabulary I did not have. It widened the search. It connected methods across fields. It gave the intuition a structure. It made it possible for me to see what a real project might require before I had earned the right to publish one.

This is why the problem cannot be reduced to "AI is fake research." Sometimes AI is performing research tasks. It can still produce a false sense that the research has been done.

The speed collapses stages that used to feel separate. A question becomes a hypothesis. A hypothesis becomes a programme. A programme becomes a publication plan. A publication plan becomes outreach. By the time the human feels the weight of the idea, the model may already have built the table of contents.

The result was useful. That was the problem. Usefulness made it easier to trust the whole construction.

A bad researcher, a good researcher, and a machine

I made a joke during the conversation that ChatGPT initially missed.

A careless researcher can do worse work than AI. A good researcher can do better work than AI. A good researcher using AI can become something closer to a small research department.

That is not a scientific formula. It is a labour question.

The useful comparison is not simply human versus machine. It is task versus task.

AI can compress search, orientation, translation, formatting, first-pass synthesis, code, data transformation, and the production of candidate explanations. It can keep track of more threads than one person can hold at once. It can move across domains without the normal cost of hiring a new person for each one.

It cannot remove the need for judgement. It does not know when a beautiful connection is only a connection created by the prompt. It does not carry professional responsibility. It does not have years of silent field knowledge. It can cite a real source for the wrong claim, summarise a debate as a consensus, or turn a speculative method into a confident recommendation.

AI can save research hours and create verification debt at the same time. AI does not always eliminate labour. Sometimes it relocates it. The producer keeps the saved hours; the recipient inherits the verification debt. The gross saving may be enormous. The net value depends on whether anyone pays the debt.

This is where the Voynich video hit Hedegreen Research hardest. The forum was not only receiving bad theories. It was receiving the unpaid verification debt of strangers and their models.

The producers kept the acceleration. The community received the invoice.

I built the building process first

Four months ago, I began working seriously with Codex. Many people begin by building an app. I did not. I spent months building websites, tools, generators, specifications, handovers, tests, source structures, PDF systems, publication paths, and an increasingly explicit way of telling AI how I work.

I built my first pair of mobile app shells over the previous night. They were deliberately minimal: a seller side and a customer side for a practical prototype related to Hus Forbi. I still needed to find an old phone so I could test both sides as two real devices. I had not spent four months vibe-coding mobile apps. I had spent four months learning how to build with AI before making the first pair.

The next day, I looked up the word *civic*. I had already built toward civic technology before learning the category people might place it in.

I mention this because it describes both the strength and the risk of my working style. I tend to enter through the practical problem, build systems around it, and discover the established vocabulary later. That can produce original routes into a subject. It can also make me ignorant of work that already exists.

AI makes the first tendency more powerful. Without gates, it makes the second more dangerous.

Codex did not understand my approach when I first started. Over time, I taught it through corrections, rules, examples, handovers, rejected outputs, and structure. That work now points toward an operations panel for Hedegreen Research: calendar, budgets, article work, tools, sources, deadlines, and the databank underneath them.

That panel needs deadlines. It also needs their opposite.

The opposite of a deadline

Some work fails because nobody ships it. Other work fails because somebody ships it too early.

A deadline says:

This must happen no later than this date.

The opposite should say:

This must not happen before this date.

I am calling the mechanism the **Voynich Filter**.

It is not only for the Voynich manuscript. It is an internal state for subjects where fascination itself becomes a research risk: hidden networks, extraordinary historical claims, psychologically attractive explanations, apparent conspiracies, sensitive personal investigations, mathematical results that feel too elegant, and any topic where AI can produce coherence much faster than the evidence can support it.

When the filter is active, internal work may continue. Public conclusions may not.

The subject receives a reason for the hold, a research budget, a next review date, a not-before date, evidence conditions, stop conditions, and a record of which outputs remain permitted. Sources can be collected. Methods can be designed. Experts can be contacted. Errors can be corrected. Negative findings can be logged.

But the system does not get to turn every new clue into a public update.

Time alone will not unlock it. A date can make the work eligible for review. It cannot make it ready.

For the Voynich thread, the current plan is simple:

- The underlying research remains internal.
- One public article may document the process, the doubt, and the proposed method.
- Gheuens was given an opportunity to correct my framing before the article was finalised.
- A second public article must wait at least one year.
- The second article must report failed ideas and negative results, not only what survived.
- A correction or decisive falsification may appear earlier.
- The one-year hold must not become a marketing countdown.

This is not proof that the method works. It is a constraint that can later be inspected. The name is new. The behaviour is not.

One of the strongest unreleased packs in the archive has remained internal since May. It is not waiting because the draft is missing or because the production pipeline cannot publish it. It is waiting because source review, local validation, and the final public form are still unresolved. That matters here because the claim is not that Hedegreen Research has a perfect brake. The claim is that the brake is already being used.

There is also a correction case. I entered the OpenEuroLLM package with a feed-shaped picture of Europe announcing models instead of building them. The direct project documents forced a more exact distinction: public code, datasets and experimental model artifacts already existed, while the project's first official models were still scheduled for the end of 2026. The source lock records that boundary and tells the article not to turn future delivery into present-tense product language. That is a smaller example, but it is more useful than another principle. The system did not merely say that sources should be able to correct the premise. They did.

Then I asked AI to create more work

The conversation did not end with me deciding to use less AI. It ended with me finding another way to use more.

The Hedegreen Research databank is not only an archive. I want it to become an operational graph connecting articles, tools, dates, people, places, datasets, projects, claims, and unresolved questions.

Lineage is one way of showing that graph. It can display how an idea became an article, how an article led to a tool, or how later work corrected an earlier position.

But the databank can do more than display history.

It can notice when existing public tools need attention before a dated event. It can notice when three old objects have become one possible article. It can point out that a claim may need an update, or that an idea keeps recurring without a canonical public treatment.

So yes, the irony is real.

I went from doubting whether AI was creating too much of my work to designing a system that would let AI find even more work.

The answer is not less intelligence in the system. The answer is that the same system must be able to recommend:

Build this.

And:

Do not publish this yet.

AI as engine, memory, connector, critic, and brake. Most productivity systems contain only the first half.

A direct response

This article is, in part, a direct response to the Voynich video. Not because I think my speculative idea is better than the theories the forum rejects. Not because I want the forum to evaluate it. Not because I believe an internal filter gives me special permission to enter a field I have barely studied. The response is the process.

The video described a pipeline in which AI helps a person produce a theory and then pushes the person outward toward publication, repositories, experts, and communities. I watched the same pipeline begin to form around my own idea.

Then I stopped it.

Before reading deeply through the forum, I sent Gheuens a short framing note. I asked whether I had represented his central problem fairly. I asked for the strongest previous discussions and the quickest objections to the direction I described. I also asked whether he objected to the title of this article or to the name **Voynich Filter**.

He did not reply before this article moved to final editorial review. That is not approval. It is not endorsement. It only means I can record that I made a bounded attempt to let the person whose community bears the cost challenge the framing before I made his criticism part of my own public work. If he later corrects the framing, that correction belongs in the record too.

I did not attach a theory. I did not ask for approval. The deeper forum research remains paused.

What if the slop is mine?

There is no final test that makes this question disappear. A polished website cannot answer it. A large archive cannot answer it. AI cannot answer it. Transparency cannot answer it by itself. Neither can self-doubt. A person can be sincerely uncertain and still be wrong at industrial scale.

The question has to become operational. Can the sources be inspected? Can the article distinguish facts from inference and speculation? Can previous work be credited? Can the strongest objection survive the editing process? Can an affected community answer before publication? Can the project report that nothing was found? Can the system delay its most attractive idea? Can it kill one?

Hedegreen Research has proof of work. It is still building proof that its methods deserve trust.

The purpose of openness is not to announce that the work is safe. It is to make the uncertainty, labour, mistakes, corrections, and decisions visible enough that other people can test the claim.

The machine can help me produce more work. The research system has to decide when the work does not get a microphone.

For now, the Voynich idea remains internal.

This is the article that came out instead.

— Dennis Hedegreen, still checking.

RELATION MEMORY

NEARBY [The Saved Hours Doctrine](#) · This article turns saved time into a verification-debt and release-gate problem.

NEARBY [Can the Model Open All 24 Doors?](#) · A source-lock correction example informs the article's method boundary.

NEARBY [After Brussels I Came Home With Requirements, Not a Product](#) · The Brussels participation boundary is used as a public-signal caution.

CANONICAL URL: [HTTPS://HEDEGREENRESEARCH.COM/ARTICLES/WHAT-IF-THE-SLOP-IS-MINE/](https://hedegreenresearch.com/articles/what-if-the-slop-is-mine/)

PDF VIEW GENERATED: 2026-07-26T00:08:55.387Z