

AI Disclosure

SCOTT Scanner (Structural and Coherent Order in Topological Transforms)

Scott Cundill

Independent Researcher, Okinawa, Japan
`scott.cam@gmail.com`

June 2026

This document describes how AI was used to develop, test and run the SCOTT Scanner. It is neither disclaimer nor defence. It is a written record so that anyone reviewing this research can understand exactly which parts involved AI and which did not.

1 Background

I am not a mathematician or physicist. I hold a BCom majoring in Information Management. I have spent thirty years building and managing software companies: I know how to scope a problem, manage an implementation, test whether something works, and ship it. I work with AI tools the same way I work with specialist contractors and development teams: I direct the work, I make the decisions, and I take responsibility for the output.

The SCOTT Scanner is an image analysis instrument that uses topological, fractal, and geometric metrics to measure spatial structure in two-dimensional images. It has been applied to human faces, AI-generated faces, coral reef tissue, quasicrystalline materials and photographic bokeh. Two papers are in submission and a series of technical notes are in progress. None of this work would exist in its current form without AI tools. None of it would exist at all without the decisions, the direction, and the judgment calls that I made throughout.

Both of those statements are true, and neither diminishes the other.

2 Which AI tools were used

Claude (Anthropic) was the primary AI partner across the project. Claude was used for code development, statistical analysis, experimental design support, some manuscript drafting and editing, structuring confound analyses, and creating handover documents for session continuity.

GPT-4o and GPT-4 (OpenAI) were used for specific tasks. Example 1: GPT-4o Vision was used as a crop-selection tool in the Bokeh probe, identifying dense regions of bokeh highlights in photographs and returning crop coordinates. GPT-4 is used mostly for image analysis to speed up the categorisation process, but I manually source and check all samples personally. Example 2: In the Bokeh probe, GPT-4o Vision was used to identify and crop regions of bokeh highlights from source photographs. The prompts were different for each category (solid vs donut bokeh), tuned to the visual features of each type. Every crop was verified by eye before inclusion.

Gemini (Google) the Pro version was used for initial code development, cross-checking analysis approaches, and manuscript feedback.

ChatGPT (OpenAI) and **Grok (xAI)** were used in code development and data analysis at minor stages. Grok was primarily used for ranking and scoring document sections to assist in their robustness.

No single model was used exclusively. Different models were used for different tasks, and in some cases the same task was checked across multiple models to catch errors.

3 What AI did

Code development. The SCOTT Scanner is written in Python. The core analysis pipeline, the metric computations, the surrogate generation, the statistical testing scripts, the cropping tools, and the reproduction scripts were all developed with significant AI assistance. In practice this means I described what I needed, the AI wrote the code, I ran it, and we iterated many times until it worked correctly, taking notes as we went. I do not write Python fluently. The AI does.

Statistical analysis. Effect sizes, confidence intervals, t -tests, confound checks, correlation analyses, and power considerations were computed with AI assistance. The AI explained the methods, wrote the computation scripts, and helped interpret the outputs. I decided which tests to run and what the results meant.

Manuscript drafting. The text of the papers and technical notes was drafted collaboratively. I provided the structure, the argument, and the voice. The AI produced drafts that I reviewed, edited, and revised. Final decisions about what to claim, what to caveat, and what to cut were mine.

Experimental design support. The design of probes, the structure of confound checks, the choice of control conditions, and the framing of predictions were developed in conversation with AI. The AI suggested approaches, flagged weaknesses, and proposed alternatives. I chose which path to take after deliberation to ensure it was in line with the overall vision.

Handover documents. The SCOTT Scanner is developed across many AI chat sessions, each with its own context window. To maintain continuity, detailed handover documents are written at the end of each session and loaded at the start of the next. These documents record the current state of every finding, every known issue, every decision made and why. The AI helps draft them, but I verify their accuracy and completeness. Managing AI is my key strength and the handover system is the governance backbone of the project.

4 What AI did not do

The predictions are mine. Predictions in this project were written by me before tests were run. That said, I learn and consult with AI constantly in my day to day continuous education, asking questions about the extremely complex degree of math as I go so I can better understand it. As a novice, I cannot purport to fully grasp the detail, which is why I chose this project in the first place: it helps me learn more about this area of science.

The judgment calls are mine. When a finding looked weak, I decided whether to pursue it or drop it. When a confound appeared, I decided how to test it. When a result contradicted expectations, I decided what it meant. I chose to drop things that were not working, often against the AI's recommendations, whose instinct is to keep probing. Examples: Finding 9 (`tau_ratio_score`) was formally killed when the data demanded it. Category B was dropped from the Bokeh replication because it was not needed for the core finding. The decision to split the project and earn our way to the next phase was a governance decision about how to manage the research honestly. These decisions were mine.

Specimen selection and verification are mine. Every image used in the project was hand selected or approved by me. Example: In the Bokeh probe, I searched Flickr manually, evaluated each candidate, downloaded it, and verified the crop. AI selected the crop region, but I checked every one.

The scientific direction and vision are mine. The idea that a topological instrument might read three-dimensional features from two-dimensional projections, the central question

behind the Bokeh probe and the broader project, is mine. AI helped me test it. AI did not conceive it.

External scientific engagement is mine. The relationship with anyone outside of this project, such as Prof. Tsutomu Ishimasa, who reviewed the quasicrystal findings and provided his own data for analysis, was initiated and maintained by me. All scientific relationship is strictly between two people, never between a person and a model.

5 Quality control and error handling

AI makes mistakes. Code has bugs, statistical reasoning can be flawed, and drafts can contain claims that exceed the evidence. The project handles this in several ways.

Cross-model checking. Where a result mattered, the analysis is checked by different AI models. Errors caught this way were corrected and documented. This is strictly part of my ability to manage AI, which has a tendency to drop things and make errors.

Handover discipline. The handover documents serve as a running audit trail. When a new session begins, the AI reads the full project state and can flag inconsistencies with its own prior outputs. This has caught multiple errors.

Running the code. Every script was run on my machine, on my data. No local agents are used: the AI cannot see my screen or my files unless I share them. When a script fails, I report the error, the AI diagnoses it, and we fix it together. The results come from actual execution, not from AI generating plausible-looking outputs.

Caveats are mandatory. Every finding in the project carries explicit caveats, written into the handover documents and into any publication. The AI helps draft them, but the rule that caveats must be stated honestly, not softened or buried, is mine. Example: When the brightness confound appeared in the Bokeh probe at $d = 0.82$, it was reported as a large effect and tested directly, not explained away.

6 The difficulty of managing AI at scale

There is a common assumption that using AI makes research easier. On small, self-contained tasks, that is true. On a project of this size, across months of development, dozens of sessions, multiple models, and a growing body of interconnected findings, managing AI is genuinely difficult, and it required every year of my thirty years of experience to keep it on track.

AI has no persistent memory between sessions. Every conversation starts from zero unless you bring the context yourself. The handover system described in Section 3 exists because without it, the project would have collapsed into contradictions within the first few weeks. Writing those documents, checking them, maintaining them, and loading the right ones at the start of each session is management overhead that the AI cannot do for you. It is project management, applied to a tool that forgets everything overnight.

AI is confidently wrong on a regular basis. It will produce code that looks correct and runs without errors, but gives the wrong answer. It will cite papers that do not exist. It will present a statistical method fluently while misunderstanding the assumption it depends on. It will agree with you when it should push back, and push back when it should agree. Catching these failures requires knowing what the right answer should look like, even when you cannot derive it yourself, and that judgment comes from experience, not from the tool.

AI drifts. Over a long session, or across a series of sessions on the same problem, AI will gradually shift its framing, its assumptions, and its level of caution. A finding that was correctly caveated in session five will be stated as confirmed in session eight, unless the handover documents hold the line. The models optimise for helpfulness, and helpfulness, left unchecked, becomes overconfidence. Managing this drift, pulling the conversation back to what is actually established rather than what the model has convinced itself is established, is constant work.

AI does not manage scope. It will happily build infrastructure for a problem you have not yet earned the right to investigate, add features nobody asked for, and pursue interesting tangents that have nothing to do with the question at hand. The governance decisions in this project, splitting into two projects rather than three, holding the Astronomy project in reserve until a finding earns it, dropping Category B from the Bokeh replication, were all decisions I made against the natural tendency of AI to expand scope. Saying no to an AI that is enthusiastically building the wrong thing is a management skill, and it is the same skill I have used for three decades with human development teams.

None of this is a complaint about AI. These are the properties of the tool, understood through experience, and managed accordingly. The point is that managing AI at scale on a project with real scientific stakes is not a shortcut. It is a different kind of work, and it is hard.

7 The philosophical position

There is a growing anxiety in academic publishing about AI-assisted research, and most of it is reasonable. The concern is that AI makes it easy to produce work that looks rigorous but is not, that the polish outpaces the substance, that claims are generated rather than earned.

This project takes a different position. AI is a tool, and like any tool, the quality of the output depends on the quality of the direction. A skilled contractor with a bad brief produces bad work. A good brief, clear requirements, honest testing, and the willingness to kill results that do not hold up, produces good work regardless of who or what writes the code.

The test is not whether AI was involved. The test is the robustness of the science: whether the results were reported honestly including failures, whether the confounds were identified and tested rather than argued away, and whether the full method is available for anyone to reproduce. Those standards apply to all research, AI-assisted or not.

This document exists because I believe transparency about the use of AI is a feature of good research, not an embarrassment. If the science is honest, the tools used to produce it are a matter of record, not a matter of concern.

Scott Cundill

Okinawa, Japan

June 2026

Repository: <https://github.com/scottcundill/scott-framework>