

Responsible Use of Generative AI in HetSys

Policy and guidance for HetSys PhD researchers and staff

EPSRC CDT in Modelling of Heterogeneous Systems, University of Warwick · Version 1.0, August 2026

This policy adapts the DTU Physics recommendations on generative AI (December 2025) and the Vancouver Guidelines (ICMJE) for the HetSys community, and extends them to the code-centred and agentic workflows widely used in computational modelling. It was informed by a community poll at the July 2026 Summer Conference, in which HetSys students and staff drew the lines of acceptable use in essentially the same places. It sits under, and does not replace, University of Warwick, funder and publisher policy.

1. Principles

Generative AI (large language models and related tools) offers real opportunities for research. In work where reliability and reproducibility are paramount, its responsible use is a matter of research integrity. Three principles hold regardless of the tool or the task, and everything else in this policy follows from them.

1. **All content is your responsibility.** If a passage of your thesis or paper turns out to be copied from an existing work, you are responsible for the plagiarism however the passage was produced. If a result is wrong, it is your responsibility. AI use never transfers responsibility away from you.
2. **AI cannot be an author, and cannot be cited as one.** An AI model cannot be accountable for the accuracy, integrity or originality of the work, so it cannot be an author, and citing an AI does not discharge principle 1.
3. **You must declare substantive use.** Declare your use of AI in publications and your thesis, following the Vancouver Guidelines and the relevant funder and publisher requirements. What must be declared will change over time; Section 6 below sets out how HetSys expects declaration to work, and why declaring is safe.

2. Colleague analogy

When you are unsure whether AI usage is acceptable, ask whether the same assistance would be acceptable coming from a helpful colleague, your supervisor, or another student.

Would it be acceptable for another student to write the introductory chapter of your thesis from a few directions you gave them? Hardly, so do not do it with AI. Would it be acceptable for a co-student to help you with some Python code for a figure? Of course, provided you check the code yourself and can stand behind it, so you may do that with AI too. The community poll at the summer conference in 2026 confirmed that this analogy is the line most researchers draw in practice, which is why it is the foundation of this policy.

The distinction the analogy captures is verifiability and responsibility. You may use AI output only where you can check it and take responsibility for it. Section 3 sets out what "checking" requires, because it means something different for code, for a literature search, for a factual claim, and for data analysis.

3. What "checking" means

"I checked it" must have an agreed meaning. Adequate checking is not the same task in each case, and in some cases checking the output alone is not enough.

Mode of use	What adequate checking requires
Code	Read it, run it, and test it. Confirm it does what is claimed, ideally with a test that would fail if it did not. For scientific code this includes physical correctness (conservation laws, limits, units), which auto-generated tests don't guarantee on their own.
Literature search	Confirm every returned reference exists and says what is claimed. This is necessary but not sufficient, since the characteristic failure is omission of relevant papers, which no amount of checking the output can reveal. Treat AI literature search as a starting point, not a complete survey, and still follow citations manually too.
Factual or conceptual claims	Corroborate against an authoritative source. Judging that an explanation "sounds right" is not checking; plausible, authoritative-sounding output can be wrong, incomplete or biased.
Data analysis and figures	Verify the procedure that produced the number or plot, not just the picture. You must be able to see, understand and rerun the steps. See the clarification in Section 4.
Agentic / project-scale code	Review generated code. Confirm tests exist and encode correctness, check the provenance of any generated analysis, and be alert to code that appears verifiable but is not. See Scenario 8.

4. Worked scenarios

The following scenarios may be useful to help decide when generative AI usage is acceptable and when it is not. In each case the verdict follows from the principles above.

1. Correcting language. Acceptable. You use AI to correct the language of text you wrote. This does not by itself need to be reported (see Section 6), though you should confirm the meaning is unchanged.

2. Literature search. Acceptable. You use AI to find literature. You remain responsible that every reference is correct and relevant, and that you have not missed relevant work (see Section 3).

3. Pasting AI text into your thesis. Not acceptable. You copy a paragraph from a chatbot directly into your thesis. You cannot be sure of the origin of the content, it may be plagiarism, and it fails the colleague analogy.

4. Curve fit you cannot inspect. Not acceptable. You paste data into a chatbot and it returns a fitted plot, with no code you can read or rerun. You cannot see how the fit was obtained, so you cannot verify it or take responsibility for it.

5. Code that produces the fit, which you check. Acceptable. You use AI to write Python code that produces a fit to your data. You read the code, confirm it does what is intended, and can rerun it, so you can take responsibility for the result.

6. Inspiration from a dialogue, checked. Acceptable. You ask an AI where 2D materials are used in industry and take inspiration for your thesis introduction, carefully checking every resulting reference and claim. You could have had the same conversation with a colleague.

7. Advice you evaluate. Acceptable. You ask an AI what to consider for convergence of an electronic-structure calculation; it raises plane-wave cut-off, k-point sampling, semi-core states. You

decide which points apply, check convergence yourself, and discuss it in your paper. You take responsibility for what to check.

8. Agentic coding across a project. Acceptable with conditions. You use an agentic tool (for example Claude Code, Codex or Gemini CLI) that writes code, runs it and iterates across your whole project, not just a single snippet, to build or refactor an analysis or simulation pipeline.

Around 40% of HetSys students already use agentic workflows (based on 2026 summer conference responses). While this is not an expectation, it is where use is growing fastest, so the policy addresses these workflows directly. The principle is unchanged, you are responsible for every line of code with your name on it, but following this looks different when an agent acts across a repository. Acceptable use requires: (i) working spec-first, deciding and validating the algorithm yourself before the agent implements it; (ii) reviewing the code the agent produces and the written summary of its work with scepticism, since both can be convincing and wrong; (iii) ensure tests exist and encode correctness, including physical correctness, with numerical tolerances chosen by the user not the AI; (iv) be able to rerun the pipeline and reproduce the result; and (v) never push AI-written code to a shared repository, or into a paper or thesis, without having run and understood it. Beware of "verification theatre" where code appears verifiable but is not. The agent has no notion of physical plausibility, so a simulation with the wrong parameterisation will run happily and give meaningless results. Checking for this is our job, not the tool's.

5. Confidential material and choice of tool

- **Confidentiality:** never paste confidential material into a standard public AI tool without the agreement of your co-authors or collaborators. This includes unpublished manuscripts, industry-partner data, and proposals or papers you are reviewing. Note that UKRI prohibits the use of AI in reviewing proposals, and most publishers prohibit it for peer review.
- **Warwick tools:** for anything involving University of Warwick information, use Microsoft Copilot or Nebula ONE, which is covered by a University data-security agreement. For generic tasks that involve no University or confidential information, for example generating generic code, you may use another AI assistant provided you hold the correct licence.
- **Dependence:** these are commercial tools that can change, restrict access or lose your data. Keep your own copies of anything you depend on, and do not build a workflow you cannot reconstruct without a particular product.

6. Declaration of AI usage

Declaration protects you and the integrity of the work; it is the research equivalent of acknowledging a colleague's help, not a confession. HetSys treats it that way, and asks you to as well. The community poll found that many students who use AI in acceptable ways would nonetheless not declare it, from a fear that declaration marks their work as lesser. That fear, rather than any intent to deceive, is the main threat to compliance.

- **What to declare:** follow the Vancouver Guidelines. Writing assistance is declared in the acknowledgements; use of AI in data collection, analysis or figure generation is described in the

methods in enough detail to allow replication, including the tool, version and prompts where relevant; AI is never listed or cited as an author.

- **How to declare:** a short standard declaration form is provided (Appendix B). Declared use that is within this policy carries no penalty and no prejudice in assessment. Supervisors are asked to declare their own AI use in co-authored outputs, so that declaring is a shared norm rather than a mark against students.
- **What need not be declared:** routine language correction (Scenario 1) need not be declared. This follows emerging practice and reflects our community's view; we will revisit it as norms settle.

7. Assessed work and taught modules

The scenarios above concern research outputs. Assessed work use is governed by the rules of the module concerned. HetSys will align module-level rules with this research policy where possible so that you face one coherent framework rather than conflicting ones. Where a module sets a stricter rule for a specific assessment, that rule applies for that assessment.

8. Training, and the value of working without AI

The strongest concern raised by our community is not about rules but about skills: that heavy reliance on AI erodes the ability to code, to reason and to notice when something is subtly wrong, precisely the judgement needed to supervise these tools.

- **Build intuition first:** in the first year or two or your PhD, do the core foundational work by hand. The friction is educational: it builds the mental model of failure modes you will later need. It is fine to use AI for boilerplate (plotting, file I/O, CI workflows, docstrings) from the start, with your supervisor's agreement.
- **Then leverage, carefully:** when you are the source of novelty, spec-first working becomes essential and some mechanical work can be delegated to AI with care. We suggest you keep one regular problem that you solve without AI, to maintain the skill.
- **Training:** HetSys training will cover the verification practices in Section 3, how to use AI as a tutor without acquiring an illusory understanding, and when to work without it. This turns "you must be able to take responsibility" into a teachable skill.

9. Status and review

This is version 1.0. Both usage and norms are moving quickly, so the policy will be reviewed annually, informed by the community poll re-run from time to time. If in doubt, apply the principles in Section 1 and the colleague analogy in Section 2, and discuss anything that remains unclear with your supervisor, exactly as you would any other question of research ethics.

Appendix A. The Vancouver Guidelines on AI, in brief

The ICMJE ("Vancouver") Recommendations, updated April 2025, are the reference standard this policy follows. The points relevant to AI are, in summary:

- authors must disclose at submission whether AI-assisted technologies were used, and describe how.
- writing assistance is reported in the acknowledgements; use in data collection, analysis or figure generation is described in the methods in enough detail to allow replication, including tool, version and prompts where applicable.
- chatbots cannot be authors, because they cannot take responsibility for accuracy, integrity and originality; humans are responsible for all submitted material, including any produced with AI.
- authors must review and edit AI output carefully, ensure there is no plagiarism, and attribute all quoted material; AI-generated material is not acceptable as a primary reference.

Consult the current ICMJE Recommendations directly for the authoritative wording, as the guidance is periodically updated.

Appendix B. Declaration of generative AI use

To accompany a thesis submission, or to adapt for a publication's cover letter and acknowledgements.

Name:	
Title of thesis / paper:	
Date:	
Signature:	

1. Following the Vancouver Guidelines, describe in the appropriate parts of the work how you used AI.
2. Provide a short overview of your overall use of AI-assisted technologies below (tools, versions, and the tasks they were used for):

--