

Annex 17. Do's and Don'ts - Artificial Intelligence in the ECDC Fellowship Training Programme

Working document version 01

The ECDC Fellowship Programme is a two-year competency-based training with two paths: the field epidemiology path (EPIET) and the public health microbiology path (EUPHEM). Both paths provide training and practical experience using a “learning by doing” approach in acknowledged training sites across European Union (EU) and European Economic Area (EEA) member states. While Generative Artificial Intelligence (Gen AI) and large language model tools can enhance learning experience and help address public health challenges, we must establish clear boundaries for Gen AI use in this educational context. Fellows must therefore critically evaluate whether Gen AI is appropriate for each task. For some assignments and module activities, you will be explicitly asked not to use Gen AI tools. Developing fundamental epidemiological skills requires dealing with challenges through your own critical thinking and problem-solving abilities.

Overreliance on Gen AI can prevent full comprehension of essential concepts. The process of encountering obstacles, working through them, and developing your own solutions is crucial for building the core competencies required of epidemiologists and public health microbiologists.

These do's and don'ts provide practical guidance and support for navigating various programme tasks while developing your professional capabilities. We expect all fellows to respect these guidelines when Gen AI use is restricted.

Do's	Don't
Follow institutional policies: always comply with your institution's AI use policies. Seek clarification if unsure. Ensure data privacy: use only publicly available or non-sensitive data (eg that does not risk identifying individuals or disclose internal information). Maintain transparency: disclose when AI tools are used in reports, presentations, or publications. Critically evaluate AI outputs: check AI-generated content for accuracy, bias, and relevance. Always validate information. Use AI for efficiency: leverage AI to streamline tasks like drafting or summarizing but retain personal judgment. Attribute sources: acknowledge AI assistance and cite sources properly. Promote ethical use: use AI in alignment with ethical standards like fairness and non-discrimination. Disable data sharing for AI model training (this prevents the data from being exposed in public domains).	Don't use AI for decisions: avoid relying, solely, on AI for decisions that impact public health or research. Don't share sensitive data: never input confidential, personal, or proprietary data into AI tools. Don't assume accuracy: AI outputs can be biased or incorrect. Don't overuse AI: avoid excessive reliance on AI, especially for critical or creative tasks. Don't violate copyright: ensure AI-generated content doesn't infringe on intellectual property. Don't use AI in restricted tasks: refrain from using AI in areas like drafting legal documents or binding agreements. Don't ignore limitations: AI can't replace human expertise or nuanced understanding. Don't use AI in defined teaching assignments: refrain from using AI when instructed in teaching context.

Guiding principles

This section is derived from the EU harmonised rules on artificial intelligence and amending Regulations (EC)

Human oversight: Ensure that Gen AI systems remain tools to support human decisions rather than replacing them. Human agency and oversight are key to maintaining ethical Gen AI use (EU AI Act, Article 14). Treat Gen AI as a collaborator, not a replacement (you steer the direction).

Data protection: Comply with data protection regulations, ensuring that Gen AI systems do not compromise the privacy of individuals (EU AI Act, Article 10).

Transparency: Make Gen AI system capabilities, limitations, and decision-making processes clear to users. Transparency fosters trust and accountability (EU AI Act, Article 13).

Risk management: Identify and mitigate risks associated with Gen AI use, particularly in high-risk applications. Ensure robust safeguards are in place (EU AI Act, Article 9).

Non-discrimination: Avoid biases in Gen AI outputs by using diverse and representative training data. Ensure Gen AI systems promote fairness and equality (EU AI Act, Article 5).

Accountability: Assign clear responsibility for the deployment and outcomes of Gen AI systems. Ensure operators are accountable for their actions (EU AI Act, Article 29).

Robustness and security: Design Gen AI systems to be resilient against misuse and errors, ensuring safe operation even in adverse conditions (EU AI Act, Article 15).

Source: Link to the Regulation (EU) 2024/1689 (EU AI Act); [Regulation - EU - 2024/1689 - EN - EUR-Lex](#) (Accessed 10 September 2025)

Table 1 represents a summary of the utility of Gen AI in relation to ethics and pedagogy.

		Pedagogy	
		Enhancement	Diminishment
Ethics	High	Revise the statistical coding (eg R) Revise the english structure of text Acknowledge Gen AI when used	Create R code from scratch Find solutions for questions in modules
	Low	NA	Create text from scratch Create fabricated data Generate list of references without revision Find solutions for modules and tasks Running data analysis in Gen AI Exclusion of stakeholders

* Definition of ethics is around the risk level for sharing personal data.

Examples of practices to avoid when using Gen AI (Table 1)

Attributing authorship to Gen AI: Gen AI tools cannot be listed as authors on scientific papers. Authorship implies accountability, which Gen AI cannot assume.

Undisclosed Use of Gen AI in manuscript writing: If Gen AI is used to write or edit scientific content, this must be transparently disclosed. Failure to do so may be considered unethical or even fraudulent.

Fabricating or manipulating data: Using Gen AI to generate synthetic data without clear labelling, explanation, and justification is considered unethical. Fabrication or manipulation of data using Gen AI tools is prohibited.

Using Gen AI-generated references without verification: AI tools can fabricate citations or provide inaccurate references. Users must verify all references and not rely solely on AI-generated bibliographies.

Failing to address bias and errors: Users are responsible for identifying and mitigating biases and errors introduced by AI models. Ignoring these issues can lead to flawed or unethical research outcomes.

Using Gen AI without stakeholder engagement: Especially in sensitive fields (e.g., health, social sciences), using Gen AI without consulting affected communities or stakeholders can be ethically problematic.

Cases for Generative Artificial Intelligence Do's and Don'ts

This document outlines practical cases for using Gen AI (e.g. ChatGPT, Gemini, Perplexity, DeepSeek, Microsoft Copilot) in everyday epidemiological tasks. It highlights cases of appropriate and inappropriate use, with examples, advantages, and risks. The goal is to guide fellows on effective, secure, and ethical use of the Gen AI.

General Principles:

- There are differences in the quality and nuance of Gen AI answers between free and paid accounts, so be mindful that outputs may depend on your subscription status. Generative AI tools carry a high risk of producing unpredictable outcomes (due to hallucination and inaccuracy), and factual errors [\[1\]](#), therefore they must be used with caution.
- Never upload sensitive, personal, or identifiable data.
- Always review and critically assess AI-generated output.
- Use AI as a supplement, not a substitute, for core skills.

Area of use	Pros / Cons	Cons
Do's (case description and example)		
Data privacy Deactivate data sharing. For ChatGPT, go to settings > Data controls > Improve the model for everyone OFF	Prevents user data from being used for model training; enhances data security	No contribution to model training

Coding Create figures in R. For example, to generate a ggplot with specific coloring, facets and annotations	- Saves time - Unnecessary to remember specific commands - The focus becomes on the analysis output and its interpretation - Allows for better visualizations	- Fellows will not understand code as well which may limit debugging abilities - Fellows with limited R experience may struggle to assess if the code is malfunctioning - Generative AI can make errors in coding due to the so-called "logic errors" and "mathematical errors" [2].
R Coding Debug code. For example, by copy-pasting the error message from R into the AI	- Saves time - Reduces need for external help - Learning opportunity	- Fellows don't practice reading help pages - Fellows don't develop independent coding skills without relying on Generative AI - Risk of hallucinations and non-functional solution
R coding Create data-cleaning script. E.g., by recoding date formats to ensure consistency and enable further analysis.	- Saves time - May reduce need for external R support - Learning opportunity	- A potential missed opportunity to fully understand what each operation is doing, leading to risk of data issues - The user should not upload data to the AI
Text writing Improve language. For example, the fellow may draft a text in basic language (without sensitive information) and use the Gen AI to improve the linguistic aspects.	- Saves time - Improves the quality of outputs - Learning opportunity to improve professional writing	- Risk of over-reliance on Generative AI - Must ensure no sensitive content is shared
Research / Learning Complement literature searches. E.g. by using the "deep research" tool to receive a summary of recent articles in a research area.	- Fast way to discover similar articles and get an overview of a new area - May identify articles not detected through search-queries	- Should not replace systematic Pubmed searches - Generative AI tend to have literature search bias, which leads to selective review for the literature [3]
Research / Learning Explain an article. For example, when reading an article with complex methodology, ask the Gen AI to break it down into simpler language or answer specific questions.	- Makes difficult articles more accessible - Less reliance of external support	- Risk of hallucinations; responses must be critically assessed - Gen AI can't confirm the facts stated in the article, i.e. the integrity of the article can't be validated by Generative AI ²
Generating research ideas To use Generative AI to generate ideas, i.e. research questions, and objectives	- Generative AI can assist the process to generate ideas	Generative AI tend to be inferior in creativity in comparison to humans' generated ideas [4]
Reviewing Review non-sensitive text for feedback. For example, to use the Gen AI review a covering letter, public health communication or professional correspondence and provide feedback on the content and structure.	- May improve the quality of first drafts of non-sensitive texts and thus lower burden on supervisors - Serves as a learning opportunity	- Sensitive information should not be shared with Generative AI if data security cannot be ensured - Generative AI does not identify mistakes with facts in the text, therefore the fellow needs to confirm the facts are correct in the text - Generative AI tends to generate false statements (known as "hallucination") - Generative AI tends to generate false references

Don'ts		
Writing Let the Gen AI independently generate text for manuscripts	- Saves time	- Risk of false information and hallucination - Lack of transparency and accountability - Generative AI can't identify false information in a text (for example when there is a name for a false microorganism, the Gen AI would treat it as a true name), - Risk of copyrights infringing if the produced text overlaps with any published material
Reviewing Share sensitive information with the Gen AI. For example, by uploading a manuscript with sensitive data for review.	- Saves time	- Risk of data security breaches
R Coding Share sensitive data. E.g. by uploading raw data to the AI with personal identification numbers	- Gen AI tools can perform fast analysis of raw data	- Data security risk - May violate GDPR or national data protection laws

^[1] Gapminder; [AI Worldview Benchmark](https://www.gapminder.org/ai/worldview_benchmark/); https://www.gapminder.org/ai/worldview_benchmark/

^[2] Sun, Y., Sheng, D., Zhou, Z. *et al.* AI hallucination: towards a comprehensive classification of distorted information in artificial intelligence-generated content. *Humanit Soc Sci Commun* **11**, 1278 (2024)

^[3] Schryen, G., Marrone, M. & Yang, J. Exploring the scope of generative AI in literature review development. *Electron Markets* **35**, 13 (2025).

^[4] Koivisto, M., Grassini, S. Best humans still outperform artificial intelligence in a creative divergent thinking task. *Sci Rep* **13**, 13601 (2023)