



NATIONAL LAW UNIVERSITY, JODHPUR

NH-65, NAGPUR ROAD
MANDORE, JODHPUR (RAJASTHAN)

PHONE 0291-2577530, 2577138

E-Mail: nlu-jod-rj@nic.in

Web Site: www.nlujodhpur.ac.in

NLU/JODH/ 2026/ 7669 -7672

Date:- 29/07/2026

Notification

The University AI Policy is hereby notified and circulated for implementation with immediate effect. The policy shall serve as the University's official framework governing the ethical, responsible, transparent, and effective use of Artificial Intelligence in teaching, learning, research, assessment, administration, and other academic activities.

All faculty staff and students are advised to carefully read the attached policy document and ensure strict compliance with its provisions. The policy shall be applicable to all members of the University community from the date of this notification.

The University AI Policy is enclosed herewith for information, reference, and necessary compliance.


Dr. Snita Pankaj, RAS
Registrar

Enclosure:

1. AI Policy.

Copy to:

1. PS to Hon'ble Vice Chancellor (for kind information)
2. All faculty, Staff and Students (through e-mail)
3. Concerned file

R
Approved,
NLSU
22-07-2026

FOR INTERNAL CIRCULATION ONLY

July 22, 2026

NATIONAL LAW UNIVERSITY JODHPUR

AI POLICY

1. Introduction and Scope

National Law University, Jodhpur (*hereinafter* referred to as 'the University') encourages informed, ethical, and responsible use of Generative AI tools. Generative AI is a rapidly evolving technology that is set to significantly transform and impact the future of work in both industry and academia. While the AI-enabled world presents unprecedented opportunities, there are also serious challenges and potential risks associated with the inherent limitations of these tools, which must be carefully considered in all the activities involving their use.

This AI Policy governs the use of AI tools and applications by the faculty and students of National Law University Jodhpur for their academic purposes. The policy guidelines are supplementary in nature and should be read in conjunction with the University's other policies in force from time to time, including, but not limited to University's Service Rule Book, Student Handbook, IP Policy and other guidelines related to IQAC and academic administration.

2. Definitions

Defined below are some common technical terms (in alphabetical order) used in context of this Policy:

Algorithm refers to a sequence of instructions for processing input data to recognize patterns, make predictions, and generate outputs.

AI agent refers to an autonomous AI program which can perform tasks and accomplish goals on behalf of a user or another system without human intervention, by designing its

1
Niranjan

own workflow. AI agents can encompass a wide range of functions beyond natural language processing including decision-making, problem-solving, interacting with external environments and performing actions.

Artificial Intelligence (AI) Alignment refers to algorithms that are used to prevent AI models (like LLMs) from generating harmful, biased, toxic/incorrect content or revealing private information (that might have crept in the training data). These techniques aim to align AI models with human values, e.g., helpful, honest, and harmless.

Bias refers to a feature of a ML model, where it tends to generate an output/prediction that unfairly favors or discriminates against certain groups or individuals based on various attributes including (but not restricted to) race, gender, or socioeconomic status. For example, a model might be biased against certain groups when approving loan applications.

Confidential Information means any information or research result belonging to the university and its stakeholders, that is not publicly known or that has been provided or received under an obligation to maintain the information as confidential.

Diffusion Models refer to a type of Generative AI that can generate realistic images (including realistic human faces that do not exist in reality) and videos. These use machine learning techniques and massive corpus of images to learn to generate new images. Some current examples of such systems include Stability AI's Stable Diffusion, OpenAI's DALL-E (beginning with DALL-E-2), Midjourney, Google's Imagen, and others.

Generative AI refers to a type of artificial intelligence which can be used to generate new content (like text, code, images, videos or music) based on its training data set. The AI uses machine learning algorithms to analyze large data sets. Two prominent technologies for GenAI are Large Language Models (LLMs) and Diffusion Models.

Hallucination refers to AI outputs (usually generated by GenAI techniques) which may prima-facie appear to be true but are in fact inaccurate or fabricated.

Large language models (LLMs) refer to a type of Generative AI that can generate human-like text in response to a prompt. They use deep learning techniques and massive corpuses of text to learn to generate a response. Some current examples of such GenAI systems include OpenAI's ChatGPT, Anthropic's Claude, Meta AI, Google's Gemini, Microsoft's Copilot, Perplexity AI, and others.

Shiraza

Machine Learning (ML), a type of Artificial Intelligence, refers to various algorithms that teach a machine/computer to learn from experiences. This is different from traditional software codes solely based on logic (set of explicit rules) and conventional algorithms. The learning happens via training data, wherein ML model figures out patterns in data and tries to generalize (learns rules implicitly) from those patterns (based on certain feedback) similar to how humans learn.

Personal Data refers to all information, whether standalone or in combination with other information, relating to an identified or identifiable individual.

Prompts are the inputs or queries that a user provides to a Generative AI application to fetch the required output. These prompts may in turn be used by the application to further train the LLMs or Diffusion Models.

Training Data refers to the content (like image, text, code, etc.) used to teach and train a machine learning system to identify patterns and learn from them and subsequently use it to make predictions.

3. Inherent limitations and potential risks associated with Gen AI

In the context of this policy, readers are expected to be mindful of the serious risks associated with Gen AI tools, as mentioned below:

1. *Biases*: Gen AI (or ML models) may produce decisions that are biased and discriminatory. These systems can retain biases present in the data used for training, which may arise during data generation and collection, dataset development, or data labeling. Furthermore, the algorithms used to train ML models may themselves introduce inherent biases due to the data, intended purpose, or evaluation methods. Biases may also emerge from differences in context, as well as from variations in interpretation or implementation during AI deployment. Consequently, the AI-generated outputs can be inconsistent with University policies and violate applicable laws.
2. *Hallucinations*: LLMs are known to generate non-existing citations, factual errors, incomplete quotes and inaccurate findings. They can also plagiarise from other sources and claim the output as 'original'.

Niraj

3. *Misinformation:* LLMs are susceptible to the classical problem of 'garbage in, garbage out' due to their dependence on the training data set. Additionally, LLMs may not have access to information and events after a certain cutoff date as applicable to their respective versions. The user can face liability for using such AI outputs which are factually incorrect.
4. *Privacy and Intellectual Property (IP) violations:* Gen AI, when given access to personal information, confidential information, copyrighted material, or trade secrets, may not be able to protect privacy rights or IP rights. The information may be revealed to a third party during their independent interaction with the tool. Despite AI alignment, AI models could still reveal private information when prompted in a particular way. Irresponsible use of Gen AI tools may thus violate University's IP Policy.
5. *Security breaches:* AI tools can be used to generate phishing emails or facilitate other cyberattacks. Any unauthorised installation and deployment of AI applications or AI agents can lead to serious security vulnerabilities.

4. Guidelines for use of Gen AI in Teaching

For teaching, learning, and pedagogy, the University believes that standardised guidelines with a one-size-fits-all approach across all courses is impractical. The University does not recommend a blanket ban on Gen AI tools. AI tools may be used by faculty and students for brainstorming, idea generation, idea exploration, and interactive online searches using LLM-enhanced search engines.

Use of Gen AI remains prohibited in *End Term Exams*, *Mid Term Exams*, and *Continuous Assessment (CA)* tests.

In courses where the instructor decides to permit the use of AI tools for *varied forms of assessment*, specific guidelines on AI usage should be developed by the instructor with the approval of Dean (Academics). These guidelines must be discussed in class by the instructor and communicated to the students enrolled in the course, ensuring that they are fully aware of the expectations during the learning process and when submitting assignments. The instructor should clearly state that these assignment-specific guidelines apply only to the particular course. For such assignments, a mandatory component of oral discussion, class presentation, or viva voce should be included as part of the evaluation, and the instructor should share the final grading criteria with the students. Students are advised to keep in mind all potential risks associated with using Gen AI tools as outlined

in Section 3 of this document, titled *Inherent limitations and potential risks associated with Gen AI*.

AI detection tools, including Turnitin, should not be the sole basis for determining policy violation or academic misconduct.

Due to the inherent limitations and risks associated with use of AI tools, as outlined in Section 3 of this document, titled *Inherent limitations and potential risks associated with Gen AI*, the use of AI tools by the faculty members for the purpose of grading of answer scripts, project reports, research papers, etc. is not allowed. Use of AI tools by faculty members for preparing written teaching material, e-content, questions papers, etc. is not allowed. The use of AI tools by instructors to support innovation in pedagogy may be explored, subject to prior approval from the Dean (Academics).

Standard tools used for copy-editing (to improve spelling, formatting, and grammar) remain outside the parameters of this guidance.

5. Guidelines for use of Gen AI in Research

The principles of ethical research are overriding and place limits to use of AI in research. Authors must ensure they remain fully accountable for any sources used and the work produced. Authors must take full responsibility for the originality, validity, and integrity of the content of their submissions.

AI tools are not allowed to be listed as coauthors on submissions. AI tools cannot be cited as authors. It is not allowed to use AI tools to create or alter images; generate text, data, or code; perform data or result analysis; develop hypotheses or conclusions. However, standard tools used for copy-editing (to improve spelling, formatting, and grammar) remain outside the parameters of this guidance.

Researchers must also adhere to the AI-related guidelines issued by the respective journals, conference organising committees, funding agencies, and other professional bodies through which they are reporting or undertaking their research work.

Gen AI tools must not be used for peer review by reviewers. Gen AI tools must not be used for generating review reports or decision letters by editors. No part of the manuscript, associated files or images, revision letter, or decision letter should be uploaded to an AI

Hiya⁵

tool, even if the purpose is copy-editing. Additionally, researchers, reviewers, and editors are expected to remain mindful of the potential risks outlined in Section 3 of this document, titled *'Inherent limitations and potential risks associated with Gen AI'*.

6. Policy Review and Amendment

This AI Policy takes effect upon approval and remains in force until amended or replaced. Given the rapidly evolving nature of this technology, the University will continue to monitor developments and gather feedback from all stakeholders to update the policy from time to time. All the feedback and suggestions for improvement may be emailed to Dean (Academics), NLU Jodhpur. Changes in this policy shall be made upon the recommendation of the AI Policy Committee and with further approval of the Hon'ble Vice Chancellor.

7. Saving Clause

The decision of the Hon'ble Vice-Chancellor shall be final on interpretation of any of the guidelines contained in this policy. The Vice-Chancellor will have the power to remove any difficulties arising in implementation of this policy.

Niraj