
Beware! The AI Act Can Also Apply to Your AI Research Practices

Anonymous Author(s)

Affiliation

Address

email

Abstract

The EU has become one of the vanguards in regulating the digital age. A particularly important regulation in the Artificial Intelligence (AI) domain is the 2024 established EU AI Act. The AI Act specifies — due to a risk-based approach — various obligations for providers of AI systems. These obligations, for example, include a cascade of documentation and compliance measures, which represent a potential obstacle to science. But do these obligations also apply to AI researchers? This position paper argues that, indeed, the AI Act’s obligations could apply in many more cases than the AI community is aware of. In our analysis of the AI Act and its applicability, we contribute the following: 1.) We give a high-level introduction to the AI Act aimed at non-legal AI research scientists. 2.) We explain with everyday research examples why the AI Act applies to research. 3.) We analyse the exceptions of the AI Act’s applicability and state that especially scientific research exceptions fail to account for current AI research practices. 4.) We propose changes to the AI Act to provide more legal certainty for AI researchers and give two recommendations for AI researchers to reduce the risk of not complying with the AI Act. We see our paper as a starting point for a discussion between policymakers, legal scholars, and AI researchers to avoid unintended side effects of the AI Act on research.

1 Introduction

Discriminatory and harmful uses of AI have caught regulators’ attention [1, 2], with the EU passing a pioneering Act on Artificial Intelligence [3], which on the one hand, seeks to protect EU citizens from the risks of AI, and on the other hand fosters the adoption and development of trustworthy AI on the EU market. Unfortunately, despite its commitment not to interfere with scientific research (Rec. 25), the AI Act has not been drafted to account for the AI research practice of publishing alongside scientific papers on platforms like GitHub, Hugging Face or Colab. As a consequence, AI researchers who publish their research artifacts at conferences like ICML, NeurIPS, or ICLR may be held liable under the AI Act.

Why should the AI and ML community care about the applicability of the AI Act to their research? As we explain below, applying the AI Act to research may result in significant compliance obligations for researchers, which they lack the time, budget and expertise to comply with. Researchers would be required to adhere to the same regulations that apply to companies like Meta, Google, or Anthropic. These requirements can create obstacles for researchers. The obstacles exist regardless of whether researchers are working within the EU or outside of it, such as in the USA. Similarly to the General Data Protection Regulation, the AI Act has an extraterritorial scope of application (Art. 2). Due to this worldwide applicability of the AI Act and the resulting obligations, it is, in our view, crucial for researchers to consider its applicability to their research practices. Where research practice

37 is not protected by the AI Act’s research exceptions, the researchers must follow all compliance
38 regulations of the AI Act. Failing to follow the AI Act’s obligations is not a trivial matter, as fines for
39 non-compliance with rules on AI up to 35,000,00€ (Art. 99 (3)).

40 Different aspects of AI regulations in context of the AI Act have already been studied, such as
41 non-discrimination regulations [4–7], deepfake regulations [8–10], the impact on fundamental rights
42 [11, 12], research ethics [13] or lobbying efforts [14]. Few studies analyse the AI Act’s or its earlier
43 drafts’ research exceptions in detail [15–20]. The Act’s research exceptions are usually discussed in
44 connection with another topic, such as copyright or data protection [21–23] or their relevance for a
45 specific application area or type of AI [24–26]. Even the most in-depth study on the AI Act’s research
46 exceptions [16] gives limited attention to the problem of regular scientific research publication
47 practices triggering the application of the AI Act. **Our position is that the AI Act’s obligations**
48 **apply in many more cases than the AI community is aware of.** This leaves researchers in uncharted
49 waters and facing severe legal challenges and liabilities. Our contribution to these challenges:

- 50 • We will give a high-level introduction of the AI Act in 2, targeting Machine Learning
51 researchers with no prior knowledge of the AI Act, and show the severe implications of AI
52 Act applicability. Furthermore, we will provide examples of AI research that fall into the
53 categories of AI systems regulated by the AI Act.
- 54 • We will analyze in-depth the exceptions of the AI Act in favour of scientific research,
55 product-oriented research, and open source. Furthermore, we explain why, in many cases,
56 the AI Act does not account for AI research practices in Section 3.
- 57 • Finally, we will reflect on legislative measures and interpretations that could mitigate the
58 challenges the AI Act poses on ML research, see Section 4.

59 2 Research and AI Act

60 We will first provide a high-level introduction to make the implications of the AI Act understandable
61 to all researchers.

62 2.1 A Primer on the AI Act

63 The Artificial Intelligence Act (AI Act) of the European Union came into force on 9 July 2024.
64 The AI Act creates a uniform legal framework for the development, placing on the market, putting
65 into service, and the use of AI systems in the EU (Rec. 1). The AI Act seeks to promote the
66 uptake of human-centric and trustworthy AI, ensuring a high level of protection of health, safety,
67 and fundamental rights, and against the harmful effects of AI while supporting innovation, including
68 science and research & development (R&D) (Art. 1, Rec. 25). To fulfil its regulatory objectives, the
69 Act sets numerous compliance obligations on the providers of the AI systems that should cover the
70 lifetime of the system (Recs. 69, 71). This raises the question of whether AI researchers also need to
71 comply with the AI Act.

72 To this end, the researcher must first assess whether they are dealing with an AI system or an AI
73 model regulated by the AI Act, and whether, through their actions, they qualify as providers of these.
74 We will cover both aspects next.

75 2.2 Categories and Processing of AI systems

76 The AI Act poses compliance obligations to to high-risk AI systems and general-purpose AI models.
77 Unfortunately, differentiating between what is an AI system and what is an AI model is challenging
78 and is subject to legal uncertainty. This is because the AI Act defines the AI system very broadly
79 (Art. (3)(1)) and in a manner that does not correspond to the language adopted in the AI research
80 community [27]. We will discuss this challenge in more detail in the Appendix A.1 and get back to
81 the issue in section 4. What is essential is that combining AI models with other components, such
82 as user interface, qualifies them as AI systems (Recital 97; [27]) regulated by the AI Act. Similarly,
83 offering demos of their research output (for example, on a website) can be seen as an AI system.

Three lines of regulations in the AI Act: prohibited AI systems, high-risk AI systems, and GPAI models. To be on the safe side, AI researchers must evaluate whether their research concerns AI categories regulated by the AI Act: prohibited AI systems, high-risk AI systems, or GPAI models ¹.

Firstly, research may involve practices that are prohibited by the AI Act, such as certain subliminal techniques, systems for exploiting vulnerabilities of a natural person or a specific group, social scoring systems, as well as systems for emotion or biometric recognition, as well as certain facial databases (Art. 5). Examples for this type of research are for example *Labeled Faces in the Wild* [28]² or Kosinski [29] who published a paper on facial recognition technology.

Secondly, the AI Act lays out detailed, technically and administratively laborious compliance obligations for high-risk AI systems. High-risk AI systems include those that represent products or safety components to products covered by the New Legislative Framework, such as machinery and medical devices (Art. 6(1); Annex III), and systems that are deployed for (Art. 6 (2); Annex III) biometrics; critical infrastructure; educational and vocational training; employment; worker's management or access to self-employment, access to and enjoyment of essential private services and essential public services and benefits; law enforcement; migration, asylum and border control management and administration of justice and democratic processes. Examples from science can also be found for high-risk systems [30–32]. Thus, if the AI Act applies to science, current research might be classified as high-risk systems, triggering a cascade of obligations. In contrast to the example of an AI system in the high-risk area mentioned above — which might sound a bit distant to researchers — the AI Act also regulates GPAI models. As a reminder, a GPAI model is an AI model that is trained with a large amount of data using self-supervision at scale, displays significant generality, and is capable of competently performing a wide range of distinct tasks (Art. 3(63)). Examples of GPAI might include models such as GPT [33] or Gemini [34]. Additionally, more open models like LLaMA [35], DeepSeek-R1 [36], Falcon [37], or Bloom [38] might also be considered GPAI models. In contrast to high-risk systems, the regulation for GPAI models is still rather vague [5] since the EU AI Office is currently working on a General Code of Practice that specifies GPAI model regulations.

Activities that trigger compliance obligations

If AI researchers are dealing with one of the regulated systems just explained, they must review whether they are likely to engage in actions that trigger AI Acts' compliance obligations for providers over the course of their research activity. The term provider refers to an actor that develops an AI system or a model. The key actions that trigger the providers obligations are "placing on the market" and "putting into service" (Art. 3(3)).

Generally, an AI system or model covered by the AI Act must comply with it the moment it is made available on the Union Market for the first time (i.e. placed on the market) (Art. 3 (9)) [39]. The term *making available on the market* means the supply of an AI system or a GPAI model for distribution or use on the Union market in the course of a commercial activity, whether in return for payment or free of charge (Art. 3 (10)). In principle, offering the possibility to download or use an AI system over the cloud could trigger the applicability of the AI Act.

Putting into service means the supply of an AI system for first use directly to the deployer or for own use in the EU for its intended purpose (Art. 3(11)). Putting into service also covers AI systems developed for in-house use [39], within the AI research institution. It can thus be established that many usual activities of AI researchers could, in principle, qualify as putting into service or making available on the market.

2.3 Legal Burdens of the AI Act to Researchers

Where AI researchers suspect that they are working on prohibited or high-risk AI systems or GPAI models regulated by the AI Act, they must assess whether their research activity triggers compliance obligations or liabilities under the AI Act. Providers must ensure that the high-risk AI systems and GPAI models they place on the market or put into service comply with all the requirements of the AI Act (Arts. 3 (3) 16, 53, 55). We will discuss examples of these severe requirements for high-risk system and GPAI models in the following.

¹Please also note that there is a fourth special category, certain AI systems in Article 50, with specific transparency obligations (CU) which are not covered in this work.

²The most popular face-recognition database at Papers with Code <https://paperswithcode.com/datasets?task=face-recognition>

For high-risk systems, a risk management system must be implemented (Art. 9 AIA). This risk management system must be planned in advance and maintained throughout the entire life cycle. For research, this would also include maintenance after publication. According to Art. 10 AIA, data governance practices need to be implemented. This includes documentation of design choices, personal data, and bias mitigation techniques. Additionally, accuracy, robustness, and cybersecurity issues must be addressed (Art. 15 AIA). This covers technical redundancy, backup, and fail-safe plans. All requirements need to be documented throughout the entire lifespan of the system (Art. 11 AIA and Art. 12 AIA). Under certain circumstances researcher must also register the high-risk-system (Art. 71 AIA)

For GPAI models, similar requirements are established (Art. 53 AIA). However, these requirements are currently being defined by the Code of Practice (see <https://digital-strategy.ec.europa.eu/en/policies/ai-code-practice>). The current draft of the Code of Practice includes roughly 50 pages of additional requirements, for example, on transparency, copyright, and safety. While details are not yet negotiated, many of the rules focus on (internal) compliance [40] and checklists. From our perspective, the AI Act's requirements pose a substantial (documentation) burden for the ordinary ML researcher.

Non-compliance with the rules poses a risk of severe penalties. The AI Act defines in Article 99 and Article 101 fines for non-compliance with the AI Act. These penalties are primarily targeted at companies (e.g., Art. 99 (3) AIA) but can also target public entities such as universities. The amount of the fine depends on whether, for example, the ban on prohibited AI systems (Art. 99(3) AIA: up to 35,000,000 €) or specific high-risk system regulations (Art. 99 (4) AIA to 15,000,000 €) have been violated.

3 Research Exceptions for AI

It should be noted that AI research concerning unregulated forms of AI does not need to comply with the Act. However, as shown above, AI research may consist of regulated forms of AI: prohibited systems, high-risk systems, and GPAI models. Where a researcher suspects that the AI Act applies to their research, they must review whether their research activity benefits from the many exceptions of the AI Act.

3.1 A Labyrinth of Exceptions

At first glance, the AI Act's commitment to the freedom of science and its aim of not undermining R&D Activity (Rec. 25) seems convincing — after all, it contains several exceptions in favour of scientific research (Art. 2(6), Rec. 109), r&d (Arts. 2 (8), 3(63), 57-62) and open-source AI (Arts. 2 (12), 53(3)). In its Digital Strategy, European Commission considers the AI Act as one of the policies promoting AI research and innovation, highlighting that "The regulation is not applicable to any activities related to AI research, testing, and development before it is marketed or put into operation." [41]. Mantelero shares this view, stating "that the AI Act is largely not applicable to research activities". [13].

However, we argue that the research exemptions of the AI Act are not drafted to account for some AI research practices, exposing researchers to a maze of exceptions, and considerable legal uncertainty about their obligations under the AI Act. (see Figure 2).³ This system of relevant exceptions for researchers is explained in plain text in the following section.

Due to the misalignment of these exceptions with AI research practices, the act of publishing an AI system or model may trigger compliance obligations — and legal liability under the Act. Without legal measures and interpretations that create legal certainty for AI research, there is a risk that the AI Act will hinder scientific research and the practice of publishing AI systems or models alongside scientific publications.

The discrepancy between the expected effect of research exceptions and their fit with the research practices of the AI research community may have originated from the lack of consideration for AI

³Please note that the Figure and our study does not account for rules on certain systems subject to transparency obligations (Art. 50) or which concern researchers deployers of AI systems (Art. 3(4)) We do not reflect on the AI Act's exceptions for AI developed for military, defense, and security purposes (Art. 2(3)).

research practices throughout the AI Act’s legislative process. The Act’s impact on academic research was not reviewed [42], and its earlier drafts were criticized for not containing research exemptions [15, 18, 19]. The Act was refined throughout the legislative process to include the research exceptions analysed below [43].

3.2 Scientific Research

The AI Act establishes an exception for scientific research that covers AI systems or AI models, including their output, specifically developed and put into service for the sole purpose of scientific research and development (Art. 2(6); Rec. 25). However, this exception covers only the act of putting the AI system or model into service.

Considering the networked nature of today’s research practice with many stakeholders and participants [16], this exception is very narrow in scope. It covers situations where the AI system is developed for the sole purpose of in-house scientific research or the direct deployment by a research collaborator (Arts. 2 (6); 3(11)). AI systems that may be used to carry out any product-oriented research, development and testing do not fulfil the requirement of “sole purpose” (Rec. 25). In practice, it is difficult for researchers to foresee whether and when their research crosses the boundary from scientific research to research and development, especially in research partnerships that may include private companies [16].

The main question is whether publishing a model belongs to the core research (or scientific) activity, which is not part of placing it on the market or putting it into service. Uploading research artefacts to repositories like GitHub, Hugging Face or Colab may simplify further research and initiate a new development loop. Please note first that “making available on the market” does not require any payment and can be done free of charge (Art. 3(10)). Thus, just because ML research artefacts can be downloaded from repositories without any fee does not imply that they have not been placed on the market (for a discussion of the open-source exception, see below).

Uploading code and models could be seen as an integral part of scientific research. Major machine learning conferences such as NeurIPS — the primary research outlets for machine learning-related research — encourage researchers to publish their models and code⁴. It is the standard (recommended) practice to publish code and models [44] when submitting a paper within the research community.

However, uploading the models and code might not belong to the core research activity. The Oxford dictionary defines research as “systematic investigation or inquiry aimed at contributing to the knowledge of a theory, topic, etc., by careful consideration, observation, or study of a subject”⁵. Still, one can argue that uploading research artefacts wouldn’t be necessary to perform the core ML research activity. According to the wording of the Oxford dictionary, it can be argued that uploading research artefacts is not part of “systematic investigation or inquiry aimed at contributing to the knowledge”. See also Figure 1 for illustration.

The act of publishing a system, in other words, placing it on the market or making it available on the EU market, might fall out of the scope of the scientific research exception.

3.3 Product-oriented research

The AI Act has a distinct exception for product-oriented research, testing, and development activity for AI systems and models (Rec. 25). The AI Act does not apply to R&D prior to the AI system or model being placed on the market or put into service (Art. 2(8)). The exception for product-oriented research seems to protect the internal R&D activity of commercial AI providers. Similarly, R&D partnerships are not carved out of the AIA, instead, the act seeks to offer AI providers with a relative freedom to experiment in the R&D phase, under the condition that the product is compliant when they seek to commercialize it. In practice, complying with the AIA requires AI providers to start planning for AIA compliance already in the R&D phase (See [16]).

Here too, product-oriented research at large companies with research units might be considered a borderline case of science [45]. Research units at large companies like Google or Meta perform research similar to that in academic institutions. Although they sometimes have a focus on product-oriented research, it might be less clear how the product-oriented research exception should be interpreted with

⁴See <https://neurips.cc/Conferences/2025/CallForPapers>.

⁵https://www.oed.com/dictionary/research_n1?tab=meaning_and_use#25922908.

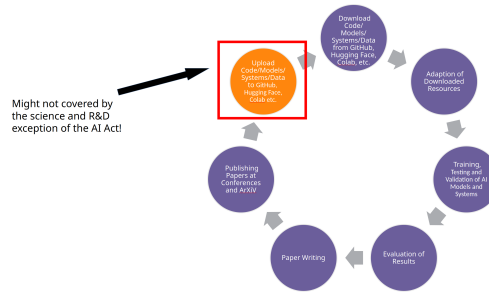


Figure 1: AI research lifecycle: Researchers start by downloading code, systems, models, and data for their research. They adapt these resources to train, test, and validate their systems and models. Afterwards, the results are evaluated. Papers are written and published at conferences and on arXiv. Simultaneously, the code, models, and data are published to public repositories. We argue that publishing code, models, and data (orange) is not always covered by the science and R&D exception of the AI Act, such that legal obligations of the AI Act might apply.

respect to work conducted at the companies’ research-oriented units. Commentators have accepted that the work conducted by units such as Google Health could, under certain circumstances, be regarded to fall under the AI Act’s exception for scientific research rather than R&D [23].

We recognize the necessity of not watering down the AI Act’s obligations for companies by overly generous interpretation of the Act’s exceptions in favour of research and product development (Art. 2(6) and 2(8)) [16, 17, 19]. However, the absence of specific rules for publishing AI systems and models in the course of research remains a serious problem. AI researchers working at companies and universities risk their research activities triggering the AI Act’s applicability for the entire R&D project involving the relevant system prematurely. This can occur also where actors in question have genuine intentions to comply with the AI Act’s regulations in the event that the AI system in question will be commercialized. This premature imposition of compliance obligations would serve neither research community nor R&D in the field of AI.

3.4 Real-Life Testing and Regulatory Sandboxes

Even when AI researchers benefit from the exceptions for product-oriented research, they must note that the exception does not extend to real-world testing of high-risk AI (Recs. 25, 141). Where the AI system is tested for its intended purpose outside a laboratory or otherwise simulated environment (Art. 3 (57)), they must either comply with the AI Act’s conditions for real-world testing such as compiling a real-world testing plan (Arts. 3 (53); 60) or conduct the testing in the context of a regulatory sandbox (Arts. 3 (55); 58-59), see Buocz et al. [46] for details.

3.5 Research Exceptions for GPAI

Due to the inconsistent usage of the terms AI system and AI model in the AI Act, the applicability of research exceptions for scientific (Art. 2(6)) and product-oriented research (Art. 2 (8)) to GPAI models is subject to further uncertainty. Both exceptions refer to AI models but do not specifically mention GPAI. Similar terminology is adopted in reference to some obligations associated with high-risk AI (Arts. 10; 15). In our view, the rest of the Act and the recitals suggest that GPAI is subject to distinct exceptions of its own.

The definition of a GPAI model includes its own product-oriented research exception. The exception applies to (GP)AI models that are used for research, development, or prototyping activities before they are placed on the market (Art. 3 (63)). The recitals affirm that compliance obligations for the providers of GPAI models do not extend to the persons developing them for scientific purposes (Rec. 109).

The AI Act’s recitals also appear to carve out a distinct research exception for GPAI models with systemic risks. The obligations for such systems should not cover GPAI models used before their placing on the market for the sole purpose of research, development, and prototyping activities (Rec.

266 97). It is notable that the research exception for GPAI with systemic risk is present only in the recitals,
267 which, unlike the text of the Act, are non-binding [47].

268 Similarly to the research exception discussed above, the research exceptions for GPAI models are
269 vague. What does “research before placing on the market” indicate? When is a GPAI model used
270 for research? Again, uploading a model to a standard ML repository like GitHub or a service like
271 Hugging Face/ Colab can be seen as part of research. When we refer to the narrow research definition
272 from the Oxford Dictionary mentioned above, one can also question this assumption (see also section
273 3.2). Therefore, the scope of the exception is, again, unclear. Furthermore, in light of the narrow
274 definition of sole scientific research for high-risk AI (Art. 2(6), Rec. 25), how can prototyping activity
275 represent the sole purpose-use of a GPAI model with a systemic risk?

276 3.6 Exceptions for Open Source AI

277 The AI Act acknowledges the role of free and open source licenses for software and data for research
278 and innovation (Rec. 102). The Act does not apply to AI systems released under free and open source
279 licenses (Art. 2(12)). Could this exception protect researchers who want to publish their AI models
280 on GitHub? It should be noted that not all models published on the platform are released under a free
281 and open source license. However, more diligent habits of relying on free and open source licensing
282 will not mitigate AI researchers’ risk of becoming subject to the AI Act’s obligations: the exception
283 for free and open source AI does not apply to high-risk AI systems (Art. 2 (12)) or GPAI models with
284 systemic risks (Art. 54 (6)). Additionally, correct licensing on GitHub is currently rather challenging
285 for the computer science community [48, 49].

286 3.7 No Exceptions for Prohibited AI

287 The AI Act’s exceptions for research and open source do not cover prohibited uses of AI (Art. 5).
288 The legal consequences for prohibited uses of AI are triggered by placing the systems on the market,
289 putting them into service, or *the use* of the said systems. Against this background, even in-house
290 research on prohibited systems could trigger liability under the AI Act.

291 3.8 Summary of Exceptions: Substantial Legal Uncertainty

292 As we have shown, at first glance, many exceptions for AI researchers exist. However, taking a closer
293 look, we showed that the exceptions might not be as broad as expected, due to their disconnect from
294 the research practices of the field of Machine Learning. Therefore, we argue that the exceptions leave
295 researchers with substantial legal uncertainty.

296 3.9 Alternative Views

297 Our viewpoint that researchers might be covered by the AI Act, can be challenged by three different
298 arguments.

299 Firstly, freedom of science is a fundamental right of the European Union (Art. 13 [50]). The AI Act
300 should not affect scientific research beyond what is specified in the Act (Rec. 25). The fact that the
301 AI Act does not refer to the publication of AI systems and models in the course of academic research
302 (Art. 2(6)) does not mean that it applies to this scientific research practice. Furthermore, scientific
303 research does not represent a commercial activity in the light of Art. 3(10).

304 We do not agree with this viewpoint since, as explained above, the AI Act directly includes research
305 exceptions and thus protects the freedom of science. Our stance is that these exceptions do not cover
306 all aspects of scientific research.

307 The second option would be to argue that publishing AI models or systems by scientific researchers
308 does not qualify as making them available to the public because this does not represent a “commercial
309 activity” (Art. 3 (10)). The Blue Guide clarifies the interpretation of EU’s product harmonization
310 legislation, which the AI Act is part of (Rec. 46). Commercial activity refers to providing goods
311 in a business-related context, the presence of which is evaluated on a case-by-case basis, taking
312 into account the regularity of supplies, the characteristics of the supplies, the characteristics of the
313 producer, the intentions of the supplies, etc. Also, non-profit organizations may be considered as
314 carrying out commercial activities if they operate in such a context. However, occasional supplies
315 by charities or hobbyists are not deemed to be part of the business context [39]. The fact that under

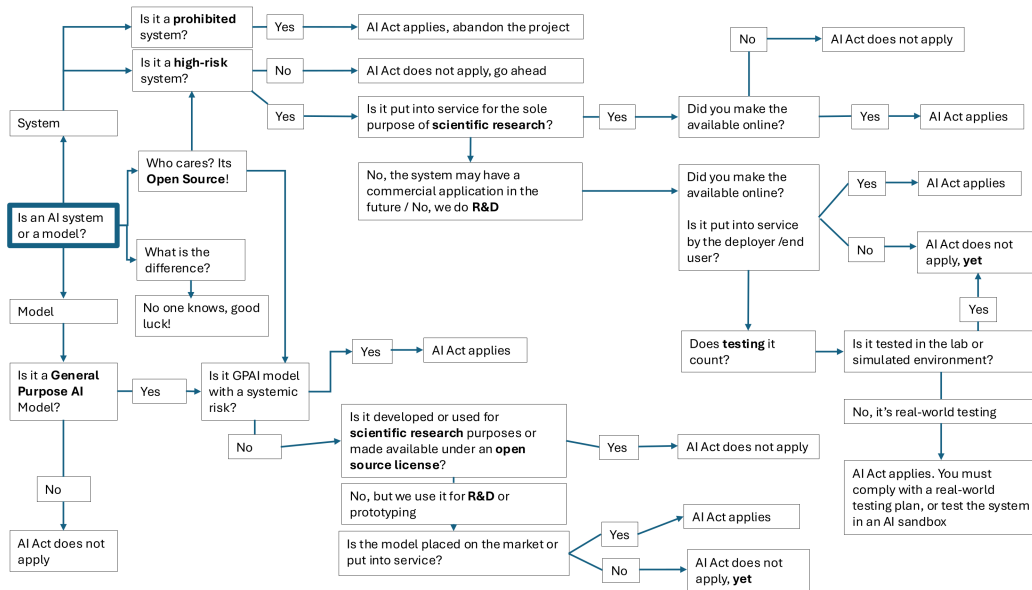


Figure 2: This flowchart illustrates the (simplified) system of AI Act’s research-relevant exceptions, which are explained in detail in section 3. It analyses whether researchers might need to comply with the AI Act as providers for prohibited and high-risk AI systems and GPAI models.

the AI Act, a provider can be a natural or legal person, public authority, agency, or other body (Art. 3(3)) suggests that a researcher could also be a provider [15]. Under the AI Act, the concept of commercial activity has a very broad meaning and also encompasses activities undertaken in a professional context. This is consistent with the Act exempting non-professional use of AI from its scope (Art. 3(4)). Furthermore, many of the high-risk uses of AI presume the use of a system by a public authority (Annex III), which could either procure a system or develop it in-house. Third, one can also argue that the EU AI Act can be seen, at least partly, as product safety law [51]. Products are outputs of industries and not of research. Thus, while the intention of the EU AI Act is to regulate products, science would be excluded from the scope. We agree that the AI Act focuses on products. Many regulations of high-risk systems and GPAI models focus on industry-related output. Nevertheless, it is important to note that the AI Act addresses technology’s impact on individuals, such as their fundamental rights [12, 46, 51, 52]. This speaks against a pure product safety regulation and instead highlights AI Act’s broader scope. Thus, scientific research is covered by the EU AI Act.

4 Recommendations: Legal Certainty is Needed for Science

What at first glance seems like a range of exceptions in favour of research is a collection of norms with complex systematics, such as differences between research exceptions in favour of high-risk AI and GPAI. The rules also contain numerous exceptions-of-exceptions. As a result, we argue that the AI Act does not align well with AI research practices and risks exposing AI researchers to compliance obligations — and liabilities. The legal uncertainty may hinder the publishing of AI models, especially on behalf of AI researchers working at commercial companies.

We hold the position that the EU should either revise the AI Act or issue an authoritative interpretation of the Act to ensure legal certainty for AI research. Recently, EU policy has displayed signs of backtracking from heavy regulation of technology to spearhead competitiveness and further investments in the R&D funding for AI [53][54]. The withdrawal of the AI Liability Directive signals a policy change [55]. We are not in favour of repealing the AIA, but position ourselves in favour of legislative drafting and revisions that account for the actual practices of AI research. In the absence of such actions, AI researchers worldwide walk on thin ice. Legal certainty should not depend on an unfortunate researcher being held liable for breaching the AI Act, and dragging them through the European judicial system to get an authoritative legal interpretation.

First, the AI Act should be revised, or an authoritative interpretation should be issued to confirm the AI researchers' freedom to publish. Furthermore, similarly to SME's and start-ups, public research institutions have "limited legal and administrative capacities" (Art. 58 (2) (g). Since research institutions contribute to the development and innovation in the field of AI, they should enjoy the same support of and consideration for the impact of compliance obligations as SMEs (recs. 109, 139, 143, 145, Art. 11 (1); 58 (2)-(3), 62 AIA) However, to protect consumer's interests, it could make sense to label these "as pre-compliant research publications" and also ensure that actors developing commercial AI systems would not use this publication route to skirt from their compliance obligations (see [16]).

In the absence of such an initiative, two interpretations could bring the AI Act in closer alignment with the AI research practices, albeit offer less legal certainty. First, the AI Act poses compliance obligations to the providers of high-risk AI systems. Consequently, the Act can be interpreted as not applying to the making available of high-risk AI models (Art. 2(6)AIA). AI researchers could attempt to publish only AI models (for example, the weights of a neural network), as opposed to AI systems (see above Section A.1 in the appendix for the difference).

The second option is to stress that the model or a system was made available for the intended purpose of scientific research. Under the AI Act, the concept of *intended purpose* refers to the use for which an AI system is intended by the provider, including the specific context and conditions of use as specified in the information supplied by the provider (Art. 3(12)). The intended purpose is determined by the provider under conditions which can be reasonably foreseen [39]. The AI system's intended purpose determines whether the system is deemed high-risk (Art. 6, Annex III). The intended purpose also qualifies the nature and scope of the provider's compliance measures (Section 2; [39]). Consequently, the provider is not liable for the consequences of a third party making significant changes either to a high-risk system (Art. 25 (1)(b)) or the *intended purpose* of a previously non-high-risk system (Art. 25 (1)(c)).

Summary of Recommendation

To reduce likelihood being qualified as providers under the AI Act, AI researchers should aim 1.) not to publish AI systems, and only models, and state that the model,⁶ 2.) is shared for the sole purpose of scientific research, and 3.) is not intended to be used for high-risk purposes, and 4.) can only be accessed and used for research purposes, and 5.) require that the person downloading it agrees that its use for any other purposes constitutes a significant change of the intended purposes, which may trigger provider's obligations and liability under the AI Act.

This interpretation would be aligned with freedom of research and the Act's objective not to interfere with scientific research (Rec. 25). Nevertheless, this interpretation does not obliterate all legal uncertainty. Firstly, the strength of contractual measures in limiting providers' liability under the AI Act would require further research. Secondly, the disclaimers should not be relied upon to escape providers' obligations for AI systems and GPAI models developed as commercializable products [16]. Thirdly, these interpretations do not apply well to GPAI. Finally, the AI Act requires providers to account for reasonably foreseeable risks and misuses. The interpretation does not hold where an AI system or a model poses obvious risks to fundamental rights, health, or safety.

5 Conclusion

AI researchers taking part in academic conferences are required to publish their AI systems or models. The act of publishing may qualify as making it available on the European market and trigger the applicability of the AI Act. Generally, AI research may concern AI categories protected by the AI Act, such as prohibited uses, high-risk AI systems, and General Purpose AI models. The AI Act's complex research exceptions do not provide a safe haven for publishing such AI systems or models. The AI Act could be interpreted to not apply to high-risk AI models and to permit AI researchers to limit their liability with disclaimers. However, the AI research community would need more legal certainty. If the EU legislator aims to control the dissemination of "forbidden knowledge" [56], the Act's research exceptions should be accompanied by safeguards fitting the process of AI research instead of exposing scientists to obligations tailored for product manufacturers. In the current state, the AI risks chilling academic research practices at the expense of scientific research and innovation as well as learning about the technology with a view of developing more trustworthy AI, thus undermining values supported by the AI Act.

⁶or where unavoidable, the AI system.

References

- [1] IAPP. Global ai law and policy tracker, 2024. URL https://iapp.org/media/pdf/resource_center/global_ai_law_policy_tracker.pdf.
- [2] White and Case. Ai watch: Global regulatory tracker, n.d. URL <https://www.whitecase.com/insight-our-thinking/ai-watch-global-regulatory-tracker>.
- [3] Regulation (EU) 2024/1689 Artificial Intelligence Act, 2024. URL <http://data.europa.eu/eli/reg/2024/1689/oj>.
- [4] Luca Deck, Astrid Schoemäcker, Timo Speith, Jakob Schöffner, Lena Kästner, and Niklas Kühl. Mapping the potential of explainable artificial intelligence (xai) for fairness along the ai lifecycle. *arXiv preprint arXiv:2404.18736*, 2024.
- [5] Kristof Meding. It’s complicated. the relationship of algorithmic fairness and non-discrimination regulations in the eu ai act. *arXiv preprint arXiv:2501.12962*, 2025.
- [6] Alessia Chiappetta. Navigating the ai frontier: European parliamentary insights on bias and regulation, preceding the ai act. *Internet Policy Review*, 12(4), 2023.
- [7] Lucía Bosoer, Marta Cantero Gamito, and Ruth Rubio-Marin. Non-discrimination and the ai act. *Law and Digitalization. Arazandi*, 2023.
- [8] Thomas Gils. A Detailed Analysis of Article 50 of the EU’s Artificial Intelligence Act, 2024.
- [9] Kristof Meding and Christoph Sorge. What constitutes a deep fake? the blurry line between legitimate processing and manipulation under the eu ai act. *arXiv preprint arXiv:2412.09961*, 2024.
- [10] Mateusz Łabuz. Deep fakes and the artificial intelligence act-an important signal or a missed opportunity? *Policy & Internet*, July 2024. Publisher: Wiley.
- [11] Alessandro Mantelero. The fundamental rights impact assessment (fria) in the ai act: Roots, legal obligations and key elements for a model template. *Computer Law & Security Review*, 54: 106020, 2024.
- [12] Isabel Kusche. Possible harms of artificial intelligence and the eu ai act: fundamental rights and risk. *Journal of Risk Research*, pages 1–14, 2024.
- [13] Alessandro Mantelero. Pre-print: The human-centric approach in scientific research: the ai act and the new frontiers of research ethics. In *Data Privacy, Data Property, and Data Sharing: An Interdisciplinary Perspective for Post-pandemic Transitional Scientific Research*. Boca Raton, CRC Press, 2025, forthcoming.
- [14] Sandra Wachter. Limitations and loopholes in the eu ai act and ai liability directives: what this means for the european union, the united states, and beyond. *Yale Journal of Law and Technology*, 26(3), 2024.
- [15] Nathalie A Smuha, Emma Ahmed-Rengers, Adam Harkens, Wenlong Li, James MacLaren, Riccardo Piselli, and Karen Yeung. How the eu can achieve legally trustworthy ai: a response to the european commission’s proposal for an artificial intelligence act. *Available at SSRN 3899991*, 2021.
- [16] Liane Colonna. The ai act’s research exemption: A mechanism for regulatory arbitrage? In *YSEC Yearbook of Socio-Economic Constitutions 2023: Law and the Governance of Artificial Intelligence*, pages 51–93. Springer, 2023.
- [17] Emre Kazim, Osman Güçlütürk, Denise Almeida, Charles Kerrigan, Elizabeth Lomas, Adriano Koshiyama, Airlie Hilliard, and Markus Trengove. Proposed eu ai act—presidency compromise text: select overview and comment on the changes to the proposed regulation. *AI and Ethics*, 3 (2):381–387, 2023.
- [18] Natali Helberger, Laurens Naudts, and Stanislaw Piasecki. The amsterdam paper: Recommendations for the technical finalisation of the regulation of gpai in the ai act, 2024.
- [19] Artur Bogucki, Alex Engler, Clément Perarnaud, and Andrea Renda. The ai act and emerging eu digital acquis. *Overlaps, gaps and*, 2022.
- [20] Katerina Yordanova. The eu ai act-balancing human rights and innovation through regulatory sandboxes and standardization, 2022.
- [21] João Pedro Quintais. Generative ai, copyright and the ai act. *Available at SSRN*, 2024.

- [22] Katherine Quezada-Tavarez, Lidia Dutkiewicz, and Noémie Krack. Voicing challenges: Gdpr and ai research. *Open Research Europe*, 2(126):126, 2022.
- [23] Katja De Vries. A researcher’s guide for using personal data and non-personal data surrogates: Synthetic data and data of deceased people. *De Lege*, 2021.
- [24] Christoph Bublitz, Fruzsina Molnár-Gábor, and Surjo R Soekadar. Implications of the novel eu ai act for neurotechnologies. *Neuron*, 112(18):3013–3016, 2024.
- [25] Sebastian Porsdam Mann, I Glenn Cohen, and Timo Minssen. The eu ai act: Implications for us health care, 2024.
- [26] Philipp Hacker, Andreas Engel, and Marco Mauer. Regulating chatgpt and other large generative ai models. In *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency*, pages 1112–1123, 2023.
- [27] Henrik Nolte, Miriam Rateike, and Michèle Finck. Robustness and cybersecurity in the eu artificial intelligence act. In *Generative AI and Law (GenLaw ’24) Workshop at 41 International Conference on Machine Learning, Vienna, Austria*, 2024.
- [28] Gary B Huang, Marwan Mattar, Tamara Berg, and Eric Learned-Miller. Labeled faces in the wild: A database for studying face recognition in unconstrained environments. In *Workshop on faces in ‘Real-Life’ Images: detection, alignment, and recognition*, 2008.
- [29] Michal Kosinski. Facial recognition technology can expose political orientation from naturalistic facial images. *Scientific reports*, 11(1):100, 2021.
- [30] Lixiang Yan, Lele Sha, Linxuan Zhao, Yuheng Li, Roberto Martinez-Maldonado, Guanliang Chen, Xinyu Li, Yueqiao Jin, and Dragan Gašević. Practical and ethical challenges of large language models in education: A systematic scoping review. *British Journal of Educational Technology*, 55(1):90–112, 2024.
- [31] Yagmur Yigit, Mohamed Amine Ferrag, Iqbal H Sarker, Leandros A Maglaras, Christos Chrysoulas, Naghmeh Moradpoor, and Helge Janicke. Critical infrastructure protection: Generative ai, challenges, and opportunities. *arXiv preprint arXiv:2405.04874*, 2024.
- [32] Prasanna Tambe, Peter Cappelli, and Valery Yakubovich. Artificial intelligence in human resources management: Challenges and a path forward. *California Management Review*, 61(4): 15–42, 2019.
- [33] Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. Language models are few-shot learners. *Advances in neural information processing systems*, 33:1877–1901, 2020.
- [34] Gemini Team, Rohan Anil, Sebastian Borgeaud, Jean-Baptiste Alayrac, Jiahui Yu, Radu Soricut, Johan Schalkwyk, Andrew M Dai, Anja Hauth, Katie Millican, et al. Gemini: a family of highly capable multimodal models. *arXiv preprint arXiv:2312.11805*, 2023.
- [35] Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, et al. The llama 3 herd of models. *arXiv preprint arXiv:2407.21783*, 2024.
- [36] Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, et al. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning. *arXiv preprint arXiv:2501.12948*, 2025.
- [37] Ebtesam Almazrouei, Hamza Alobeidli, Abdulaziz Alshamsi, Alessandro Cappelli, Ruxandra Cojocaru, Mérouane Debbah, Étienne Goffinet, Daniel Hesslow, Julien Launay, Quentin Malartic, et al. The falcon series of open language models. *arXiv preprint arXiv:2311.16867*, 2023.
- [38] Teven Le Scao, Angela Fan, Christopher Akiki, Ellie Pavlick, Suzana Ilić, Daniel Hesslow, Roman Castagné, Alexandra Sasha Luccioni, François Yvon, Matthias Gallé, et al. Bloom: A 176b-parameter open-access multilingual language model, 2023.
- [39] European Commission. *The ‘Blue Guide’ on the implementation of EU product rules 2022*, volume 65. The European Union, 2022/C.
- [40] Frederik Zuiderveen Borgesius, Philipp Hacker, Nina Baranowska, and Alessandro Fabris. Non-discrimination law in europe, a primer. introducing european non-discrimination law to non-lawyers. *Introducing European non-discrimination law to non-lawyers (April 7, 2024)*, 2024.

- [41] European Commission. European research development and deployment of ai, 2025. URL <https://digital-strategy.ec.europa.eu/en/policies/european-ai-research>.
- [42] European Commission. Commission staff working document impact assessment accompanying the proposal for a regulation of the european parliament and of the council laying down harmonised rules on artificial intelligence (artificial intelligence aact) and amending certain union legislative acts, 2021. URL <https://ec.europa.eu/newsroom/dae/redirection/document/75792>.
- [43] European Parliament and Council. Proposal for a regulation laying down harmonized rules on artificial intelligence (artificial intelligence act) and amending certain union legislative acts, four columns, 2024. URL <https://artificialintelligenceact.eu/wp-content/uploads/2024/01/AIA-Final-Draft-21-January-2024.pdf>.
- [44] Joelle Pineau, Philippe Vincent-Lamarre, Koustuv Sinha, Vincent Larivière, Alina Beygelzimer, Florence d’Alché Buc, Emily Fox, and Hugo Larochelle. Improving reproducibility in machine learning research (a report from the neurips 2019 reproducibility program). *Journal of machine learning research*, 22(164):1–20, 2021.
- [45] Asia J Biega and Michèle Finck. Reviving purpose limitation and data minimisation in data-driven systems. *arXiv preprint arXiv:2101.06203*, 2021.
- [46] Thomas Buocz, Sebastian Pfotenhauer, and Iris Eisenberger. Regulatory sandboxes in the ai act: reconciling innovation and safety? *Law, Innovation and Technology*, 15(2):357–389, 2023.
- [47] Tadas Klimas and Jurate Vaiciukaite. The law of recitals in european community legislation. *ILSA J. Int’l & Comp. L.*, 15:61, 2008.
- [48] Christopher Vendome, Gabriele Bavota, Massimiliano Di Penta, Mario Linares-Vásquez, Daniel German, and Denys Poshyvanyk. License usage and changes: a large-scale study on github. *Empirical Software Engineering*, 22:1537–1577, 2017.
- [49] Thomas Wolter, Ann Barcomb, Dirk Riehle, and Nikolay Harutyunyan. Open source license inconsistencies on github. *ACM Transactions on Software Engineering and Methodology*, 32(5):1–23, 2023.
- [50] CFREU. Charter of fundamental rights of european union, 2000. URL https://www.europarl.europa.eu/charter/pdf/text_en.pdf.
- [51] Marco Almada and Nicolas Petit. The eu ai act: a medley of product safety and fundamental rights? *Robert Schuman Centre for Advanced Studies Research Paper*, (2023/59), 2023.
- [52] Alina Wernick. Impact assessment as a legal design pattern—a “timeless way” of managing future risks? *Digital Society*, 3(2):29, 2024.
- [53] Mario Draghi. The future of european competitiveness. part a. Technical report, European Commission, 2024. URL https://commission.europa.eu/document/download/97e481fd-2dc3-412d-be4c-f152a8232961_en?filename=The%20future%20of%20European%20competitiveness%20_%20A%20competitiveness%20strategy%20for%20Europe.pdf.
- [54] European Parliament. Questionnaire to the commissioner-designate henna virkkunen executive vice-president for tech sovereignty, security and democracy, 2024. URL https://hearings.elections.europa.eu/documents/virkkunen/virkkunen_writtenquestionsandanswers_en.pdf.
- [55] Cristina Frattone. Anatomy of a fall: On the anticipated withdrawal of the ai liability directive proposal, 2025. URL <https://verfassungsblog.de/anatomy-of-a-fall-aiact-aild-pld/>.
- [56] Thilo Hagendorff. Forbidden knowledge in machine learning reflections on the limits of research and publication. *AI & Society*, 36(3):767–781, 2021.

A Appendix

A.1 Models vs. Systems in the AI Act

The AI Act differentiates between AI systems and models. The AI Act’s definition of an AI *system* is very broad, covering a vast range of theoretical and applied AI research: “machine-based system

that is designed to operate with varying levels of autonomy and that may exhibit adaptiveness after deployment, and that, for explicit or implicit objectives, infers, from the input it receives, how to generate outputs such as predictions, content, recommendations, or decisions that can influence physical or virtual environments” (Art. 3(1)).

Besides AI systems, the AI act regulates GPAI *models*, which are defined as “AI model, including where such an AI model is trained with a large amount of data using self-supervision at scale, that displays significant generality and is capable of competently performing a wide range of distinct tasks regardless of the way the model is placed on the market and that can be integrated into a variety of downstream systems or applications” (Art. 2(63)).

If the AI Act regulates high-risk AI *systems*, can an AI *model* developed by an AI researcher also fall under this definition? The challenge is that the AI Act defines AI system in a manner that does not correspond to the language adopted in the AI research community [27]⁷. This exposes researchers to considerable legal uncertainty because publishing a system is more likely to trigger liability under the AI Act than publishing an AI model. Looking, for example, at the NeurIPS Call for Papers⁸ we note that an ML model is the learned algorithmic structure itself (e.g. weights of a neural network). A Machine Learning system also includes the hardware, libraries, etc. AI Models require the addition of further components, such as, for example, a user interface, to become AI systems (Recital 97; [27]). However, even if we assume a very narrow definition of an AI model, at least, research offering demos of their research output (for example, on a website) can be seen as an AI system.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper’s contributions and scope?

Answer: [Yes]

Justification: We contribute to the debate on whether the AI Act applies to research.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [Yes]

Justification: We contribute to the debate on whether the AI Act applies to research.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.

⁷To add to the confusion, some of the AI Act’s obligations for high-risk systems refer to models that are part of those systems (Art. 15 [27]).

⁸<https://neurips.cc/Conferences/2025/CallForPapers>.

- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [NA]

Justification:

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [NA]

Justification:

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed

instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.

- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [NA]

Justification:

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [NA]

Justification:

Guidelines:

- The answer NA means that the paper does not include experiments.

- 710 • The experimental setting should be presented in the core of the paper to a level of detail
711 that is necessary to appreciate the results and make sense of them.
712 • The full details can be provided either with the code, in appendix, or as supplemental
713 material.

714 **7. Experiment statistical significance**

715 Question: Does the paper report error bars suitably and correctly defined or other appropriate
716 information about the statistical significance of the experiments?

717 Answer: [NA]

718 Justification:

719 Guidelines:

- 720 • The answer NA means that the paper does not include experiments.
721 • The authors should answer "Yes" if the results are accompanied by error bars, confi-
722 dence intervals, or statistical significance tests, at least for the experiments that support
723 the main claims of the paper.
724 • The factors of variability that the error bars are capturing should be clearly stated (for
725 example, train/test split, initialization, random drawing of some parameter, or overall
726 run with given experimental conditions).
727 • The method for calculating the error bars should be explained (closed form formula,
728 call to a library function, bootstrap, etc.)
729 • The assumptions made should be given (e.g., Normally distributed errors).
730 • It should be clear whether the error bar is the standard deviation or the standard error
731 of the mean.
732 • It is OK to report 1-sigma error bars, but one should state it. The authors should
733 preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis
734 of Normality of errors is not verified.
735 • For asymmetric distributions, the authors should be careful not to show in tables or
736 figures symmetric error bars that would yield results that are out of range (e.g. negative
737 error rates).
738 • If error bars are reported in tables or plots, The authors should explain in the text how
739 they were calculated and reference the corresponding figures or tables in the text.

740 **8. Experiments compute resources**

741 Question: For each experiment, does the paper provide sufficient information on the com-
742 puter resources (type of compute workers, memory, time of execution) needed to reproduce
743 the experiments?

744 Answer:

745 Justification: [NA]

746 Guidelines:

- 747 • The answer NA means that the paper does not include experiments.
748 • The paper should indicate the type of compute workers CPU or GPU, internal cluster,
749 or cloud provider, including relevant memory and storage.
750 • The paper should provide the amount of compute required for each of the individual
751 experimental runs as well as estimate the total compute.
752 • The paper should disclose whether the full research project required more compute
753 than the experiments reported in the paper (e.g., preliminary or failed experiments that
754 didn't make it into the paper).

755 **9. Code of ethics**

756 Question: Does the research conducted in the paper conform, in every respect, with the
757 NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

758 Answer: [Yes]

759 Justification:

760 Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes] ,

Justification: We contribute to the debate on whether the AI Act applies to research and discuss positive and negative impacts.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification:

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [NA]

Justification:

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, `paperswithcode.com/datasets` has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification:

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and research with human subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA] .

Justification:

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional review board (IRB) approvals or equivalent for research with human subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification:

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. Declaration of LLM usage

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification:

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.