
Carbon Literacy for Generative AI: Visualizing Training Emissions Through Human-Scale Equivalents

Mahveen Raza*
Independent Student Researcher
Toronto, Canada
mahveen.raza10@gmail.com

Maximus Powers†Clarkson University
Arizona, United States
maximuspowersdev@gmail.com

Abstract

Training large language models (LLMs) requires substantial energy and produces significant carbon emissions that are rarely visible to creators and users, due to a lack of transparent data available. We compile reported and estimated carbon emissions (kg CO₂) for 13 state-of-the-art models (2018–2024) during their training to reflect the environmental severity of these emissions. These carbon emissions values are translated to human-friendly equivalences, trees required for absorption and average per-capita footprints, as well as scaled comparisons across household, commercial, and industrial contexts through our interactive demo. Our key takeaways note a lack of transparency surrounding reported emissions during model training. Furthermore, the amount of emissions in only training data is alarming, causing harm that cannot be mitigated quickly enough by the environment. We position this work as a socio-technical contribution that bridges quantitative emissions analysis with human-centered interpretation to advance sustainable and transparent AI practice. By offering an accessible lens on sustainability, it promotes more responsible engagement with generative AI in creative communities. Our interactive demo is available at: <https://neurips-c02-viz.vercel.app/>.

1 Introduction

GenAI models have achieved remarkable capabilities in creating human-like text and images, but these advances come with substantial energy costs that translate into large carbon emissions [23, 20]. Training state-of-the-art large language models (LLMs) often requires hundreds or thousands of GPUs running for weeks, consuming vast amounts of electricity. For example, it is estimated that training OpenAI’s GPT-3 (175 billion parameters) consumed about 1,287 MWh of electricity and resulted in roughly 502 metric tons of CO₂ emissions [20]. This single training run’s emissions is equivalent to $\approx 20,080$ trees’ annual absorption, and to ≈ 104 years of an average human’s emissions. As GenAI moves from research labs to widespread deployment, understanding and communicating its environmental footprint becomes crucial for sustainable AI development.

Prior research in this area falls into two main strands, but both remain incomplete. The first set of studies estimates the CO₂ emissions of commercial and research LLMs, since such data are rarely disclosed in technical reports. These estimates are necessarily indirect and rely on limited information such as FLOPs, GPU hours, or hardware specifications, which introduces uncertainty. The second strand highlights that even when emissions are reported, the sheer scale, often measured in metric tons of CO₂ equivalent, can appear abstract and detached from everyday experience. Although some work proposes translating emissions into human-scale analogies (e.g., personal activities or familiar energy

*

†

usage) to improve public understanding [23], this approach has not yet been systematically applied to state-of-the-art generative models. This gap motivates our work: we compile and standardize reported and estimated emissions for 13 leading models and reframe them through transparent, human-centered visualizations.

In this work, we address this sustainability gap by compiling training emissions for state-of-the-art models, sourced from technical reports where available or estimated otherwise, and presenting them in accessible, human-scale equivalents. Our goal is to raise awareness in the research community and the public about the environmental impact of model training and to encourage the adoption of efficiency-oriented techniques. Our contributions are:

- We compile reported or estimated *training* emissions for 13 GenAI models and translate them into human-friendly comparisons such as tree absorption [8] and per-capita footprints [27].
- We provide an interactive demo that illustrates these comparisons, offering an accessible perspective on the scale of emissions and their broader environmental consequences.

2 Related Work

Carbon Footprint of AI Models. Growing awareness of AI environmental impact has led to numerous studies measuring the carbon footprint of model training. [25] sounded an early alarm by showing that training a transformer-based NLP model with hyperparameter search emitted CO₂ on the order of hundreds of kilograms to tons, comparable to cross-country flights [6]. Subsequent analyses by Patterson et al. [20] and others refined these estimates and identified key factors influencing emissions: model size, training duration, hardware efficiency, energy grid carbon intensity, and data center cooling overhead [6] [16]. For instance, training the 11-billion-parameter T5 model was shown to consume less energy when using customized TPU hardware in a high-efficiency Google data center, compared to GPUs in a standard facility [20] [13]. Such studies suggest that choice of model architectures, processors, and training location can yield 100×–1000× differences in carbon emissions for the same task [6]

In response to these challenges, the concept of Green AI has emerged, calling for energy efficiency and carbon footprint to be treated as primary evaluation metrics alongside model accuracy [24] [20]. [24] advocated for reporting the computational cost of ML experiments and favoring approaches that achieve comparable results with less resource consumption. Since then, the research community has initiated various efforts to enable “greener” AI. Techniques such as model distillation, network pruning, and sparsity have been shown to reduce training and inference costs significantly, sometimes by 50–90% with minimal performance loss [12] [10]. The emergence of more efficient model architectures (e.g. transformer variants like the Switch Transformer) and hardware accelerators (TPUs, AI chips) also contributes to cutting the energy per training operation [9]

Another positive development is the growing transparency from AI labs regarding energy use. Some organizations now publish model cards with environmental impact metrics. For instance, some organizations now publish model cards with environmental metrics, such as Meta’s model cards for LLaMA-3 [18] and LLaMA-2 [17]. Tools like *Experiment Impact Tracker* [11] and *CarbonTracker* [2] were developed to make it easier for researchers to log energy usage during training and estimate emissions based on regional electricity carbon intensity. In this work, we compile reported and estimated training emissions of state-of-the-art LLM families and reframe them through human-friendly comparisons.

3 Methodology

Model Selection (2018–2024). We analyzed 13 prominent GenAI models, from early architectures like BERT [5] and GPT-2 [21] to large-scale releases such as GPT-4 [19], the LLaMA family [26, 18], and DeepSeek v3 [4], using reported or estimated training emissions.

Emission Data Collection For each model, we gathered available training emissions data from published reports and prior sustainability analyses. Disclosed values (e.g., LLaMA-2 and LLaMA-3) are marked as *reported* (*R*), while others are *estimated* (*E*). In the absence of official data, estimates were derived from industry model cards containing FLOPs, GPU hours, or hardware specifications

[20, 1]. These known figures provided baselines for estimating comparable models using standardized methods from prior Green AI studies. It is crucial to see that each approach introduces uncertainty: FLOP-based estimates assume hardware efficiency; GPU-hour methods depend on utilization and duration; scaling methods assume similar infrastructure and efficiency across generations. Based on comparisons between reported and estimated values in earlier work [20], [16] [2], we estimate an average uncertainty margin of $\pm 25\%$, reflecting variability in datacenter efficiency, hardware utilization, and regional carbon intensity. Together, these methods offer a reasonable approximation of training emissions and their environmental impact.

To make emissions more interpretable, we translate raw CO₂ values into two main equivalents:

Tree Absorption. We estimate the number of trees required to absorb the model’s emissions using the standard assumption that one tree absorbs ≈ 25 kg of CO₂ per year [8]. The formula is:

$$\text{Trees Required} = \frac{\text{Emissions (kg)}}{25}$$

Human Equivalence. We compare model emissions to an average human’s yearly CO₂ footprint, taken as ≈ 4.8 tonnes (4800 kg) per year [27]. The equivalence is expressed as:

$$\text{Human Years} = \frac{\text{Emissions (kg)}}{4800}$$

4 Results

Table 1: Training emissions of GenAI models (2018–2024), with equivalent impacts. Tree absorption assumes 25kg CO₂/year [8], and average per-capita footprint is 4,800kg/year [28]. R = reported, E = estimated. Provenance: For each row, the Model column cites the original paper/card for dataset/training context, while the CO₂ (kg) column cites the source for the reported or estimated emissions value used in our calculations.

Model	Year	CO ₂ (t)	CO ₂ (kg)	Type	Trees (25 kg/yr)	Human yrs (4,800 kg/yr)
BERT (base) [5]	2018	0.652	652 [7]	E	26.1 trees	≈ 0.14 yrs (≈ 1.7 mo.)
BERT-Large [5]	2018	2.6	2,600 [1]	E	104 trees	≈ 0.54 yrs (≈ 6.5 mo.)
GPT-2 (OpenAI) [21]	2019	0.735	735 [7]	E	29.4 trees	≈ 0.15 yrs (≈ 1.8 mo.)
RoBERTa [15]	2019	5.5	5,500 [1]	E	220 trees	≈ 1.15 yrs
GPT-3 (175B) [3]	2020	502	502,000 [20]	E	20,080 trees	≈ 104.6 yrs
BLOOM (176B) [14]	2022	25	25,000 [20]	E	1,000 trees	≈ 5.21 yrs
OPT (175B) [28]	2022	70	70,000 [20]	E	2,800 trees	≈ 14.6 yrs
Gopher (280B) [22]	2022	352	352,000 [20]	E	14,080 trees	≈ 73.3 yrs
GPT-4 (OpenAI) [19]	2023	5,184	5,184,000 [1]	E	207,360 trees	$\approx 1,080$ yrs
LLaMA-2 (70B) [26]	2023	539	539,000 [17]	R	21,560 trees	≈ 112.3 yrs
LLaMA-3.1 (405B) [18]	2024	8,930	8,930,000 [1]	E	357,200 trees	$\approx 1,860$ yrs
LLaMA-3 (70B) [18]	2024	2,290	2,290,000 [18]	R	91,600 trees	≈ 477 yrs
DeepSeek v3 [4]	2024	597	597,000 [1]	E	23,880 trees	≈ 124 yrs

The results in Table 1 show an exponential rise in training emissions as model size increases. Early models like BERT (2018) and GPT-2 (2019) produced under 1 tonne of CO₂, while frontier models such as GPT-4 and LLaMA-3.1 reached thousands of tonnes, requiring hundreds of thousands of trees for absorption. Only LLaMA-2 and LLaMA-3 reported official emissions; all others rely on indirect estimates [20, 1], underscoring the need for standardized reporting. Human-scale equivalents make this tangible: GPT-4’s footprint equals $\approx 1,080$ years of an average human’s emissions, LLaMA-3.1 $\approx 1,850$ years, while GPT-2 corresponds to less than two months. These sharp jumps reflect how compute demands outpace linear scaling laws.

Emission Jumps. With the rise of each major AI model release, emissions have increased rapidly, compounding over time. From 2019 to 2020, emissions jumped from 5.5 t (RoBERTa) to 502 t

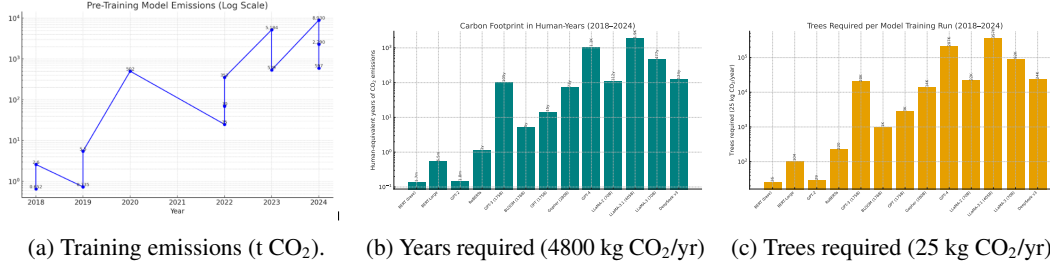


Figure 1: **Carbon Literacy Dashboard.** Three aligned views of training emissions (2018–2024) for 13 GenAI models: **(a)** scientific units (t CO₂), **(b)** human-scale impact (average human-years; 4.8 t CO₂/yr), and **(c)** ecological framing (trees needed; 25 kg CO₂/yr). Values in (b,c) are derived from (a) via: $\text{Human Years} = \frac{\text{Emissions (kg)}}{4800}$ and $\text{Trees} = \frac{\text{Emissions (kg)}}{25}$. Using a log y-axis in (a) makes early models visible while highlighting the sharp jumps (e.g., RoBERTa → GPT-3 → GPT-4).

(GPT-3), a **91× increase**. From 2020 to 2023, they rose again to 5,184 t (GPT-4), nearly 10× higher than GPT-3. Even within a single year, steep increases are visible: in 2022, emissions grew from 70 t for OPT to 352 t for Gopher, a **5× increase**. This pattern reflects several factors, including growing model size, larger datasets, and longer training runs on more powerful hardware. As models expand in parameters and training data, they require longer time and power resources, driving higher emissions. Since every new release grows in performance, carbon footprint rises sharply with each generation.

Notably, BLOOM (176B) exhibits substantially lower training emissions than models of comparable scale such as GPT-3 or OPT. This difference likely stems in training infrastructure: BLOOM was trained on A100 GPUs within energy-efficient data centers powered partly by renewable energy. This highlights how hardware choice and environmental safeguards can meaningfully reduce the carbon footprint of large-scale model training.

5 Discussion

Environmental and Social Implications. Scaling GenAI carries significant environmental costs, with training emissions of frontier models rivaling those of entire communities. LLaMA-3 alone produced $\approx 2,290$ t CO₂, equivalent to ≈ 477 average human-years of emissions raising concerns about the resulting environmental burden. These trends provide an outlook to the future of training emissions requiring even more energy as newer and larger models are released. Without stronger efficiency and mitigation, this trend will exacerbate climate risk. This provides an opportunity for developers to develop sustainable training practices in terms of GPUs, FLOPs, and hardware. However, this relies transparency from developers.

Transparency and Accountability. Of the 13 models reviewed, only LLaMA-2 and LLaMA-3 disclosed official emissions with remaining models requiring estimation from FLOPs, GPU hours, or hardware specifications, highlighting the need for standardized emissions reporting. While companies may be reluctant to disclose such data due to reputational risks, these figures remain crucial to identify and promote sustainable AI practices. Initial steps towards mitigation and transparency are being taken. Google and DeepMind operate carbon-neutral data centers and Hugging Face maintains an open Green AI Dashboard to track model energy usage. Expanding these initiatives will be key to ensuring accountability.

Assumptions and Limitations. Our estimates are approximations: tree absorption (25 kg CO₂/year) varies by species and region, and human per-capita footprints (4.8 tonnes) differ globally. We use global averages for comparability, to make model emissions accessible despite inherent uncertainty. Another limitation is the focus on solely training emissions. Inference and deployment, for example, contribute substantially to total emissions and can often exceed training emissions.

Future Directions. Sustainable AI requires transparent reporting, efficiency-focused research, and strong mitigation efforts from both developers and users, who can promote energy-efficient and environmentally conscious practices. While our study only looks at training emissions, future work

should also account for lifecycle emissions, including inference and deployment to get the full scope of the effect of emissions.

References

- [1] AI Index Steering Committee. Ai index report 2025. Technical report, Stanford Institute for Human-Centered Artificial Intelligence (HAI), 2025. URL https://hai.stanford.edu/assets/files/hai_ai_index_report_2025.pdf. Accessed: 2025-08-26.
- [2] Louise Anthony, Benjamin Kanding, and Raghavendra Selvan. Carbontracker: Tracking and predicting the carbon footprint of training deep learning models. In *Proceedings of the 37th International Conference on Machine Learning (ICML) Workshop on Tackling Climate Change with Machine Learning*, 2020.
- [3] Tom B. Brown et al. Language models are few-shot learners. *arXiv preprint arXiv:2005.14165*, 2020. URL <https://arxiv.org/abs/2005.14165>.
- [4] DeepSeek-AI. Deepseek llm: Scaling open-source language models with 10t tokens. *arXiv preprint arXiv:2405.04434*, 2024. URL <https://arxiv.org/abs/2405.04434>.
- [5] Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. Bert: Pre-training of deep bidirectional transformers for language understanding. In *Proceedings of NAACL-HLT*, 2019. URL <https://arxiv.org/abs/1810.04805>.
- [6] Jesse Dodge, Taylor Prewitt, Remi Tachet des Combes, Erika Odmark, Roy Schwartz, Emma Strubell, Alexandra Sasha Luccioni, Noah A Smith, Nicole DeCario, and Will Buchanan. Measuring the carbon intensity of AI in cloud instances. In *Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency*, pages 1877–1894, 2022.
- [7] Register Dynamics. Artificial footprints series: The environmental impact of ai, 2024. URL <https://www.register-dynamics.co.uk/blog/artificial-footprints-series-the-environmental-impact-of-ai>. Accessed: 2025-08-27.
- [8] EcoTree. How much co₂ does a tree absorb? <https://ecotree.green/en/how-much-co2-does-a-tree-absorb>, 2025. Accessed August 2025; estimates annual CO₂ absorption of 10–40 kg per tree, approximately 25 kg/year on average.
- [9] William Fedus, Barret Zoph, and Noam Shazeer. Switch transformers: Scaling to trillion parameter models with simple and efficient sparsity. In *International Conference on Learning Representations (ICLR)*, 2021. URL <https://arxiv.org/abs/2101.03961>.
- [10] Song Han, Huizi Mao, and William J. Dally. Deep compression: Compressing deep neural networks with pruning, trained quantization and huffman coding. In *International Conference on Learning Representations (ICLR)*, 2016. URL <https://arxiv.org/abs/1510.00149>.
- [11] Peter Henderson, Jieru Hu, Joshua Romoff, Emma Brunskill, Dan Jurafsky, and Joelle Pineau. Towards the systematic reporting of the energy and carbon footprints of machine learning. In *Proceedings of the International Conference on Machine Learning (ICML)*, 2020.
- [12] Geoffrey Hinton, Oriol Vinyals, and Jeff Dean. Distilling the knowledge in a neural network. *arXiv preprint arXiv:1503.02531*, 2015. URL <https://arxiv.org/abs/1503.02531>.
- [13] Norman P. Jouppi, Cliff Yoon, Morgan Ashcraft, and et al. A domain-specific supercomputer for training deep neural networks. In *Proceedings of the 2020 ACM/IEEE 47th Annual International Symposium on Computer Architecture (ISCA)*, pages 1–14, 2020. doi: 10.1109/ISCA45697.2020.00010.
- [14] Teven Le Scao, Angela Fan, Christopher B. Akiki, Ellie Pavlick, Suzana Ilić, et al. Bloom: A 176b-parameter open-access multilingual language model. *arXiv preprint arXiv:2211.05100*, 2023. URL <https://arxiv.org/abs/2211.05100>.

- [15] Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Mandar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer, and Veselin Stoyanov. Roberta: A robustly optimized bert pretraining approach. In *arXiv preprint arXiv:1907.11692*, 2019. URL <https://arxiv.org/abs/1907.11692>.
- [16] Alexandra Sasha Luccioni, Simon Viguier, and Anne Ligozat. Estimating the carbon emissions of machine learning models. *arXiv preprint arXiv:2206.05233*, 2022. URL <https://arxiv.org/abs/2206.05233>.
- [17] Meta AI. Llama 2: Open foundation and fine-tuned chat models. <https://huggingface.co/meta-llama/Llama-2-70b>, 2023. Accessed: 2025-08-27.
- [18] Meta AI. Llama 3: Advancing open foundation models. <https://huggingface.co/meta-llama/Meta-Llama-3-70B>, 2024. Accessed: 2025-08-27.
- [19] OpenAI. Gpt-4 technical report, 2023. URL <https://arxiv.org/abs/2303.08774>.
- [20] David Patterson, Joseph Gonzalez, Quoc V. Le, Chen Liang, Lluís-Miquel Munguia, Daniel Rothchild, David R. So, Maud Texier, and Jeff Dean. Carbon emissions and large neural network training. *arXiv preprint arXiv:2104.10350*, 2021. URL <https://arxiv.org/abs/2104.10350>.
- [21] Alec Radford, Jeff Wu, Rewon Child, David Luan, Dario Amodei, and Ilya Sutskever. Language models are unsupervised multitask learners. Technical report, OpenAI, San Francisco, CA, 2019. URL https://cdn.openai.com/better-language-models/language_models_are_unsupervised_multitask_learners.pdf. OpenAI Technical Report.
- [22] Jack W. Rae, Sebastian Borgeaud, Trevor Cai, Katie Millican, Jordan Hoffmann, Francis Song, et al. Scaling language models: Methods, analysis & insights from training gopher. *arXiv preprint arXiv:2112.11446*, 2021. URL <https://arxiv.org/abs/2112.11446>.
- [23] Kate Saenko. Ai’s growing carbon footprint: Why we need carbon literacy. *Scientific American*, August 2023. URL <https://www.scientificamerican.com/article/ais-growing-carbon-footprint-why-we-need-carbon-literacy/>.
- [24] Roy Schwartz, Jesse Dodge, Noah A. Smith, and Oren Etzioni. Green ai. *Communications of the ACM*, 63(12):54–63, 2019. URL <https://arxiv.org/abs/1907.10597>.
- [25] Emma Strubell, Ananya Ganesh, and Andrew McCallum. Energy and policy considerations for deep learning in nlp. In *Proceedings of ACL*, 2019. URL <https://arxiv.org/abs/1906.02243>.
- [26] Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, et al. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288*, 2023. URL <https://arxiv.org/abs/2307.09288>.
- [27] Worldometers. Co2 emissions per capita. <https://www.worldometers.info/co2-emissions/co2-emissions-per-capita>, 2025. Accessed August 2025; reports global per-capita CO₂ emissions of 4.8 tons for 2022 (and includes country-level data).
- [28] Susan Zhang, Stephen Roller, Naman Goyal, Mikel Artetxe, Moya Chen, Shuohui Chen, et al. Opt: Open pre-trained transformer language models. *arXiv preprint arXiv:2205.01068*, 2022. URL <https://arxiv.org/abs/2205.01068>.

NeurIPS Paper Checklist

1. Claims

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

Answer: [\[Yes\]](#)

Justification: The abstract and introduction clearly state the contribution: compiling reported and estimated emissions for 13 models (2018–2024), translating them into human-friendly comparisons, and highlighting transparency gaps. The claims match the presented results.

3. Limitations

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

Answer: [Yes]

Justification: The paper explicitly discusses assumptions (average tree absorption, global per-capita averages) and limitations (incomplete reporting, reliance on estimates), framing them as approximations.

5. Theory assumptions and proofs

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

Answer: [NA]

Justification: The paper does not include theoretical results or formal proofs; it is empirical and descriptive.

7. Experimental result reproducibility

8. 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: [Yes]

Justification: All emission values are either cited from official sources or estimated using documented methods (FLOPs, GPU hours, standardized formulas). Tables and formulas are provided to allow reproducibility of calculations.

9. Open access to data and code

10. 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: [Yes]

Justification: The interactive demo and dataset of compiled emissions are made available online through an anonymized demo URL that is provided in the paper, enabling public access to the data and results.

11. Experimental setting/details

12. 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: [Yes]

Justification: The methodology section details model selection, emission data collection, and equivalence calculations. No hyperparameters or training runs are involved since the work is observational.

13. Experiment statistical significance

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

Answer: [NA]

Justification: The paper does not present stochastic experiments requiring error bars or confidence intervals. It is based on reported or estimated fixed values.

15. Experiments compute resources

16. Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [NA]

Justification: The paper does not involve new model training. It compiles previously reported/estimated training emissions of existing models.

17. Code of ethics

18. Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics

Answer: [\[Yes\]](#)

Justification: The work conforms to NeurIPS ethics guidelines, as it uses publicly available reports and sustainability studies, cites all sources, and raises awareness of environmental costs.

19. Broader impacts

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

Answer: [\[Yes\]](#)

Justification: The discussion covers both positive impacts (raising awareness, promoting sustainability, informing creative practices) and negative aspects (large emissions, inequitable distribution of burdens).

21. Safeguards

22. 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: The paper does not release high-risk models or datasets, only emissions summaries and visualizations.

23. Licenses for existing assets

24. 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: [\[Yes\]](#)

Justification: All existing models, datasets, and reports are properly cited with their original sources, licenses, or official technical reports (e.g., OpenAI, Meta, Stanford AI Index).

25. New assets

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

Answer: [\[Yes\]](#)

Justification:

Justification: The paper introduces a compiled dataset and interactive demo of emissions, with documentation describing data sources, calculation methods, and assumptions, along with references for transparency and reproducibility.

27. Crowdsourcing and research with human subjects

28. 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: The paper does not involve crowdsourcing or human subject research.

29. Institutional review board (IRB) approvals

30. 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: No human subject experiments are conducted.

31. **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: LLMs were not part of the core methodology. Any LLM use was limited to writing/editing support, not central to the scientific contribution.