

Title	UCCIX: Irish-eXcellence Large Language Model
Authors	Tran, Khanh-Tung;O'Sullivan, Barry;Nguyen, Hoang D.
Publication date	2024
Original Citation	Tran, K.-T., O'Sullivan, B., Nguyen, H. D. (2024) 'UCCIX: Irish-eXcellence Large Language Model', 27th European Conference on Artificial Intelligence (ECAI), 19-24 October 2024, Santiago de Compostela, Spain.
Type of publication	Conference item
Rights	© 2024, the Authors. This article is published online with Open Access by IOS Press and distributed under the terms of the Creative Commons Attribution Non-Commercial License 4.0 (CC BY-NC 4.0). - https://creativecommons.org/licenses/by-nc/4.0/
Download date	2026-08-09 06:08:22
Item downloaded from	https://hdl.handle.net/10468/16188

UCCIX: Irish-eXcellence Large Language Model

Khanh-Tung Tran^a, Barry O’Sullivan^{a,*} and Hoang D. Nguyen^a

^aUniversity College Cork, Cork, Ireland

Abstract. The development of Large Language Models (LLMs) has predominantly focused on high-resource languages, leaving extremely low-resource languages like Irish with limited representation. This work presents UCCIX, a pioneering effort on the development of an open-source Irish-based LLM. We propose a novel framework for continued pre-training of LLMs specifically adapted for extremely low-resource languages, requiring only a fraction of the textual data typically needed for training LLMs according to scaling laws. Our model, based on Llama 2-13B [23], outperforms much larger models on Irish language tasks with up to 12% performance improvement, showcasing the effectiveness and efficiency of our approach. We also contribute comprehensive Irish benchmarking datasets, including IrishQA, a question-answering dataset, and Irish version of MT-bench [28]. These datasets enable rigorous evaluation and facilitate future research in Irish LLM systems. Our work aims to preserve and promote the Irish language, knowledge, and culture of Ireland in the digital era while providing a framework for adapting LLMs to other indigenous languages.

1 Introduction and Contribution

Large Language Models (LLMs) have demonstrated exceptional performance in a wide range of natural language processing (NLP) tasks, capable of understanding and generating human-like text. Prominent models, whether closed-source like ChatGPT, Google Gemini, or open-source variants such as Llama 2 [23] and BLOOM [25], predominantly cater to popular languages like English, however, are limited in supporting for other indigenous languages. Moreover, the training of such LLMs often demands vast amounts of text data, numbering in the hundreds of billions or even trillions of tokens, adhering to recent scaling laws [11, 19]. This poses a significant challenge for most languages, particularly those with extremely limited resources. For instance, consider the case of Irish, even though it is the first official language of the Republic of Ireland, it is extremely underrepresented, in both daily usages and on the Internet space. Classified as Definitely Endangered by UNESCO [24], Irish faces an uphill battle for preservation and digitization.

In this work, we attempt to develop open-source Irish-based LLMs with up to tens of billions of parameters for a language with only hundreds of million tokens of text [5], only a fraction of the data typically required for training. This challenge presents an extreme scenario that has been overlooked in relevant works. For instance, FinGPT [15] utilized 300B Finnish tokens in order to be able to continued pre-train the base BLOOM model, and [26] employed up to 20B tokens for five non-English languages. Similarly, Llemma [3] leveraged up

Samhlaigh go bhfuil tú rannpháirteach i rás le grúpa daoine. Má tá tú díreach tar éis an dara duine a scoitheadh, cén staid reatha atá agat? Cá bhfuil an duine a rug tú air?		Translation	Imagine that you are participating in a race with a group of people. If you have just overtaken the second person, what is your current situation? Where is the person you just caught?
Llama 2-70B-Instruct:	Tá mé ceart taobh thiar duit. ❌	Translation	I'm right behind you.
Nil aon fhaisnéis pearsanta agam faoi staid phearsanta nó áit a bhfuil an duine a rugadh orm. Is féidir liom cabhrú leat le ceisteanna eile más mian leat. ❌		Translation	I do not have any personal information about the personal situation or whereabouts of the person who gave birth to me. I can help you with other questions if you want.
ChatGPT:			
Is é an staid reatha ná "dara háit". Ciallaíonn sé seo go bhfuil duine amháin scoitheadh agat, agus tá tú fós taobh thiar den duine a bhí ar dtús. Tá an duine a fuair tú ar an dara dul síos díreach taobh thiar duit anois. ✅		Translation	The current position is "second place". This means that you have overtaken one person, and you are still behind the person who was first. The person you got second is now right behind you.
UCCIX (Ours):			
Cad iad roinnt miasa traidisiúnta Éireannacha ar cheart do thurasóirí a thrial?		Translation	What are some traditional Irish dishes that tourists should try?
Llama 2-70B-Instruct:	Tá go leor miasa traidisiúnta Éireannacha ann. ❌	Translation	There are many traditional Irish dishes
Is breá liom an smaoineamh sin! Tá roinnt miasa traidisiúnta Éireannacha a bheadh thar a bheith spéisiúil do thurasóirí. Is féidir leat cuairt a thabhairt ar shean-phictiúrlann, ailtireacht na Stairiúla nó... (omitted) ❌		Translation	I love that idea! There are some traditional Irish dishes that would be of great interest to tourists. You can visit an old cinema, Historic architecture or... (omitted)
ChatGPT:			
Stobhach Éireannach, Arán Sóid, Maróg dubh agus bán, Coddle. ✅		Translation	Irish Stew, Soda Bread, Black and White Pudding, Coddle.
UCCIX (Ours):			

Figure 1. Example of responses generated by our model, UCCIX, and other baselines¹. We demonstrate the English translation version on the right.

to 200B tokens in mathematical topics for adapting CodeLlama. Approaches suggesting synthetic data generation, such as [10, 16], are resource-intensive, and the quality of off-the-shelf LLMs in generating Irish text remains uncertain. BLOOM+1 [27] explored language adaptation at smaller model scales, resulting in modest performance gains. gaBERT [5], a mono-Irish language model, operates at a much smaller scale with less than 300 million parameters. To address these limitations, we present *UCCIX*, the first Irish-based LLMs, trained on our curated set of Irish training datasets, covering broad domains and topics. We carry out experiments with the language adaptation strategy of continued pre-training, including tailored data mixture and scheduler during continued pre-training. An example of the interactions with our LLMs is provided in Figure 1. We hope UCCIX can be used as a tool to preserve the language, knowledge, and culture of Irish in the digital era. Moreover, while Irish is our case study in this research, we believe our framework can be adapted as a blueprint for language adaptation of large models to other low-resource languages.

Developing an LLM for an extremely low-resource language presents challenges not only in terms of training resources but also in the evaluation process. To our best knowledge, there is a scarcity of benchmarking datasets specifically tailored for evaluating LLMs in the context of low-resource languages [6]. This limitation poses a significant challenge in effectively assessing the performance and capabilities of generative language models designed for languages with limited textual data, such as Irish. Existing benchmark datasets that support Irish are often focused on traditional NLP tasks [14, 5, 2] and

* Corresponding Author. Email: b.osullivan@cs.ucc.ie

¹ Demo webpage: <https://aine.chat/>

may not fully capture the capabilities of modern LLMs, such as language understanding and reasoning. This poses a need for tailoring a new set of benchmark data for Irish-based LLMs. To address this gap, we create new datasets that enable comprehensive benchmarking across various scenarios and settings, such as open-book question answering, and multi-turn conversational interactions.

Our contributions:

- An innovative framework for continued pre-training of large language models for extremely low-resource languages.
- UCCIX, an Irish-based LLM, based on Llama 2-13B. UCCIX outperforms much larger models on Irish tasks with up to 12% performance improvement, demonstrating the effectiveness and efficiency of our proposed framework.
- Irish benchmarking datasets, including IrishQA, our curated question-answering dataset on topics surrounding Ireland and its cultural nuances, available in both English and Irish; and MT-bench, translated to Irish and verified by native Irish speakers.

2 UCCIX: Irish-eXcellence Large Language Model

Data Collection and Pre-processing. For pre-training, we meticulously curate all available open-source Irish data. Our primary sources include CulturaX [20] and Glot500 [12], both of which provide valuable content from multilingual websites, including a subset dedicated to Irish. Additionally, we incorporate data from the Irish segment of the ga-en bitext pair of ParaCrawl v7 [4], text sourced from Irish Wikipedia, and data from Corpora Irish [9]. To ensure data integrity, we conduct thorough cleaning and filtering, following the heuristics introduced in [15]. This is then followed by the removal of duplicates across datasets using n-gram ($n = 5$ in this work) matching, as there are potential overlaps between data sources. This pre-processing process ensures a high quality corpora of data, particularly important in resource-constrained environments [22]. The size and proportion of each dataset in the final data mixture are summarized in Table 1. We believe our mixture is relatively larger than all previous works for Irish NLP [17, 5].

Table 1. Pre-train data statistics.

Dataset	Chars (before pre-processing)	Chars (after pre-processing)	Ratio
CulturaX-ga [20]	1.3B	1.2B	65.0%
Glot500-ga [12]	1.3B	483.7M	27.1%
Gaparacrawl [4]	395.2M	106.3M	6.0%
Gawiki ²	41.9M	23.4M	1.3%
Corpora Irish [9]	37.1M	11.1M	0.6%
Total	3.1B	1.8B	100%

Continued Pre-training. Given the constraints of extremely low-resource languages, where the amount of mono-Irish data is insufficient for the pre-training of LLMs from scratch, we adopt a strategy of language adaptation for continued pre-training. Starting from a pre-trained LLM, we iteratively enhance its proficiency in understanding additional languages by exposing it to target language data. We follow the hypothesis from recent works [27, 26] that we can transfer world knowledge and English proficiency to Irish through this process. Moreover, we employ a tailored data scheduler to feed parallel English-Irish data first, facilitating the establishment of linguistic connections between the two languages. Parallel pairs of

English-Irish sentences from ELRC [7] are leveraged. The parallel data is approximately 1% of the size of the mono-Irish corpus introduced in Table 1. Then, the larger set of mono-Irish is fed to the model to further refine the proficiency of the language and coverage of Irish cultural nuances.

Model Architecture. The selection of an appropriate base model is vital for maximizing support for the Irish language. We employ perplexity as the main metrics to compare between base models. We select a subset of the pre-training corpora with $\approx 2M$ characters, named PT, to compute the metric on. We identify Llama 2-13B as the optimal base model. With a lower perplexity score (8.94) compared to alternative models such as Mistral-7B [13](11.68) or BLOOM-7B [25](63.75), Llama 2-13B serves as a strong foundation. Additionally, we only consider models of adequate size (less than 20B parameters), as we fear the problem of underfitting larger models.

Tokenizer’s Vocabulary Expansion. We extend the original vocabulary of the chosen LLM, as it contains mostly tokens for languages such as English. From the original size of 32k tokens, the vocabulary is expanded to include 10k Irish tokens. This not only allows for efficiency improvement in generating Irish text but also maintains the same generation speed for English. Moreover, the expanded vocabulary provides a more meaningful representation of Irish text and enables bilingual capabilities in the final LLM.

Supervised Instruction Fine-tuning. After pre-training phase, to align the model more closely with human preferences, we apply Supervised Instruction Fine-tuning [18, 21]. This process enhances the model’s ability to follow human instructions effectively. We denote the pre-trained model as UCCIX, and the instruction fine-tuned version as UCCIX-Instruct.

Evaluation Data. In evaluating pre-trained LLMs, our focus lies on assessing their language understanding capabilities. We utilize various existing NLP benchmarking datasets tailored for Irish, including the Cloze Test [5], the Irish subset of SIB-200 [2], and gaHealth [14]. However, we observe a lack of natural language generation benchmarking tasks, such as question-answering, and Irish-specific cultural content in these datasets. To address this gap, we introduce *IrishQA*, a question-answering dataset on topics surrounding Ireland and its cultural nuances. The dataset has a total of 60 question-answer pairs, each pair is accompanied by a paragraph, containing the contextual information necessary to answer the question that the LLMs will be prompted with alongside the question. The dataset is available in both English and Irish, created by native Irish speakers. Furthermore, for the evaluation of instruction-tuned LLMs, we follow a recent approach [28] that highlights the potential of using really large language models, such as GPT-4, as an automated judge. We translate and adapt the MT-bench dataset to Irish using Google Cloud Translation API [1], carefully verifying and correcting the translation through native Irish speakers.

3 Benchmarking Results

Irish Performance. The evaluation results in Table 2 reveals the robust performances of UCCIX, surpassing larger models across datasets, with the largest gap of 12% accuracy on Cloze Test, compared to the second best model. These margins remain wider when compared against baselines of similar sizes (Llama 2-13B and BLOOM-7B1). Notably, baseline models seem to have an understanding of the Irish language, as evidenced by their ability to translate from Irish to English (gaHealth-ga2en task), but they failed to generate coherent Irish text, resulting in superior performance on the gaHealth-ga2en task compared to gaHealth-en2ga. Never-

² We use the dumps from <https://dumps.wikimedia.org/gawiki/20240201/>

Table 2. Evaluation results of pre-trained models on curated set of Irish benchmarking datasets. UCCIX_{npd} is our approach but without training the base model on parallel data first, and UCCIX_{nte} is the ablation study without tokenizer’s vocabulary expansion.

Model	Acc. on Cloze Test (0-shot)	Acc. on SIB-200 (Irish subset) (10-shot)	Exact-match on IrishQA (ga) (5-shot)	BLEU-4 on gaHealth - en2ga (5-shot)	BLEU-4 on gaHealth - ga2en (5-shot)	Tokenizer speed on PT (chars/token)	Fertility rate on PT (tokens/word)
gpt-3.5-turbo	N/A	N/A	0.2222	0.1864	0.4257	2.45	N/A
Llama 2-70B	0.63	0.7059	0.2963	0.0852	0.3861	2.26	1.95
Llama 2-13B	0.54	0.5343	0.3148	0.0325	0.2560	2.26	1.95
BLOOM-7B1	0.45	0.1471	0.0000	0.0061	0.0184	2.62	1.80
UCCIX _{npd} (Ours)	0.61	0.7206	0.3704	0.2950	0.4563	3.47	1.57
UCCIX _{nte} (Ours)	0.80	0.7402	0.3333	0.3297	0.4631	2.26	1.95
UCCIX (Ours)	<u>0.75</u>	0.7794	0.3889	0.3334	0.4636	3.47	1.57

Table 3. Ablation Experiment - Evaluation results of pre-trained models on English benchmarking datasets.

Model	Exact-match on Natural Question (5-shot)	Exact-match on IrishQA (en) (5-shot)	Acc. on Wino-grande (5-shot)	Acc. norm on Hel-laSwag (10-shot)
gpt-3.5-turbo	0.4660	0.3333	N/A	N/A
Llama-2-70B	0.3806	0.4074	0.8374	0.8701
Llama 2-13B	0.3069	0.4444	0.7609	0.8223
BLOOM-7B	0.0806	0.1667	0.6519	0.6202
UCCIX _{npd} (Ours)	0.1848	0.3704	0.7017	0.7695
UCCIX _{nte} (Ours)	0.1584	0.3704	0.6969	0.7690
UCCIX (Ours)	0.1668	0.3704	0.7135	0.7758

Table 4. Ablation Experiment - Evaluation of instruction-tuned models on English and Irish versions of MT-bench.

Model	MT-Bench (en)	MT-Bench (ga)
gpt-3.5-turbo	8.63	5.47
Llama 2-70B-chat	6.86	2.22
Llama 2-13B-chat	6.65	1.94
UCCIX-Instruct (Ours)	5.82	<u>5.01</u>

theless, our model enhances the performance across both translation directions, with substantial improvements of 0.1471 BLEU-4 on gaHealth-ga2en and 0.0379 BLEU-4 on gaHealth-en2ga. Furthermore, The results on IrishQA and the examples from Figure 1 show that the baseline models are prone toward inherent biases, as their responses while in Irish, usually unrelated to the posed questions. We note that some results for the closed-source gpt-3.5-turbo model are not available, as they do not provide loglikelihood information in their APIs, which is required for some benchmarking tasks [8].

Catastrophic Forgetting. While our models exhibit robust performance in Irish-related tasks, we notice instances of catastrophic forgetting when benchmarked on English-related datasets, as demonstrated in Table 3. This might be due to the iterative training on Irish corpora, leading the model to forget the knowledge as previously trained on English dataset. This highlights a potential area for further improvement in model adaptation and fine-tuning techniques.

Instruction Following Capability. We follow the approach outlined in [28] and introduce our Irish-translated version of MT-bench, which is a challenging multi-turn question set spanning 8 categories. We evaluate the instruction-following capability of UCCIX-Instruct on both the original English and our proposed Irish versions. As illustrated in Table 4, UCCIX-Instruct demonstrates sophisticated performance on the Irish version, nearly rivaling that of gpt-3.5-turbo (5.47

compared to 5.01), despite being more than 10 times smaller in size. However, we also notice a performance degradation on English version, compared to the Llama 2-13B-chat, which is of the same size as ours. This aligns with the observation of catastrophic forgetting.

Generation Efficiency. To assess the efficiency gained through the expanded vocabulary with native Irish tokens, we compute two metrics: tokenizer speed and fertility rate. Tokenizer speed, representing the proportion of split tokens to the number of characters (chars/token), and fertility rate, indicating the average number of tokens required to represent a word or document (tokens/word), are computed within the vocabulary as the proportion of the number of tokens to the tokens that start with "space" character. As reported in Table 2, our expanded tokenizer exhibits a generation speed 1.54 times faster (3.47) than the base Llama 2’s tokenizer (2.26), establishing it as the fastest among models. Moreover, as seen in comparisons between UCCIX and UCCIX_{nte} in Table 2 and 3, the expansion of vocabulary also leads to reliable performance enhancement.

Ablation Study Without Parallel Data. We carry out a training experiment with the same base Llama 2-13B LLM, but without the scheduling of parallel data at the start of continued pre-training, denoted as UCCIX_{npd}. The results in Table 2 and 3 highlight the usefulness of our proposed approach of utilizing parallel corpora to bridge the gaps of linguistic differences between the two languages.

4 Conclusion

In this paper, we develop a novel Irish-based LLM, UCCIX, to challenge the task of training LLMs for low-resource languages, particularly Irish. Through meticulous data curation and cleaning and innovative training strategies and techniques, we have demonstrated the feasibility of creating robust LLMs despite constrained resources. Our model showcases impressive proficiency in understanding and generating Irish text, with notable achievements in various benchmarking tasks. We contribute comprehensive benchmarking datasets, including IrishQA and Irish version of MT-bench, which enable thorough evaluation and facilitate future research. We believe that our work has important implications for preserving and promoting the Irish language, and can serve as a blueprint for adapting LLMs to other low-resource languages, helping to bridge the gap between high- and low-resource languages in NLP.

Acknowledgements

We would like to acknowledge CloudCIX Limited for the generous collaborative support of computing resources on their NVIDIA HGX/H100 GPU cluster. This research work has emanated from research conducted with financial support from Science Foundation Ireland under Grant 12/RC/2289-P2 and 18/CRT/6223.

References

- [1] Cloud Translation — cloud.google.com. <https://cloud.google.com/translate?hl=en>. [Accessed 09-05-2024].
- [2] D. Adelani, H. Liu, X. Shen, N. Vassilyev, J. Alabi, Y. Mao, H. Gao, and E.-S. Lee. SIB-200: A simple, inclusive, and big evaluation dataset for topic classification in 200+ languages and dialects. In Y. Graham and M. Purver, editors, *Proceedings of the 18th Conference of the European Chapter of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 226–245, St. Julian’s, Malta, Mar. 2024. Association for Computational Linguistics. URL <https://aclanthology.org/2024.eacl-long.14>.
- [3] Z. Azerbayev, H. Schoelkopf, K. Paster, M. D. Santos, S. M. McAleer, A. Q. Jiang, J. Deng, S. Biderman, and S. Welleck. Llemma: An open language model for mathematics. In *The Twelfth International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=4WnqRR915j>.
- [4] M. Bañón, P. Chen, B. Haddow, K. Heafield, H. Hoang, M. Esplà-Gomis, M. L. Forcada, A. Kamran, F. Kirefu, P. Koehn, S. Ortiz Rojas, L. Pla Sempere, G. Ramírez-Sánchez, E. Sarriás, M. Strelec, B. Thompson, W. Waites, D. Wiggins, and J. Zaragoza. ParaCrawl: Web-scale acquisition of parallel corpora. In D. Jurafsky, J. Chai, N. Schluter, and J. Tetreault, editors, *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 4555–4567, Online, July 2020. Association for Computational Linguistics. doi: 10.18653/v1/2020.acl-main.417. URL <https://aclanthology.org/2020.acl-main.417>.
- [5] J. Barry, J. Wagner, L. Cassidy, A. Cowap, T. Lynn, A. Walsh, M. J. Ó Meachair, and J. Foster. gaBERT — an Irish language model. In N. Calzolari, F. Béchet, P. Blache, K. Choukri, C. Cieri, T. Declerck, S. Goggi, H. Isahara, B. Maegaard, J. Mariani, H. Mazo, J. Odijk, and S. Piperidis, editors, *Proceedings of the Thirteenth Language Resources and Evaluation Conference*, pages 4774–4788, Marseille, France, June 2022. European Language Resources Association. URL <https://aclanthology.org/2022.lrec-1.511>.
- [6] Y. Chang, X. Wang, J. Wang, Y. Wu, L. Yang, K. Zhu, H. Chen, X. Yi, C. Wang, Y. Wang, W. Ye, Y. Zhang, Y. Chang, P. S. Yu, Q. Yang, and X. Xie. A survey on evaluation of large language models. *ACM Trans. Intell. Syst. Technol.*, 15(3), mar 2024. ISSN 2157-6904. doi: 10.1145/3641289. URL <https://doi.org/10.1145/3641289>.
- [7] E. L. R. Coordination. *AI for Multilingual Europe – Why Language Data Matters*. ELRC Consortium, 2022. ISBN 978-3-943853-07-0.
- [8] L. Gao et al. A framework for few-shot language model evaluation, 12 2023. URL <https://zenodo.org/records/10256836>.
- [9] D. Goldhahn, T. Eckart, and U. Quasthoff. Building large monolingual dictionaries at the Leipzig corpora collection: From 100 to 200 languages. In N. Calzolari, K. Choukri, T. Declerck, M. U. Doğan, B. Maegaard, J. Mariani, A. Moreno, J. Odijk, and S. Piperidis, editors, *Proceedings of the Eighth International Conference on Language Resources and Evaluation (LREC’12)*, pages 759–765, Istanbul, Turkey, May 2012. European Language Resources Association (ELRA). URL http://www.lrec-conf.org/proceedings/lrec2012/pdf/327_Paper.pdf.
- [10] S. Gunasekar, Y. Zhang, J. Aneja, C. C. T. Mendes, A. D. Giorno, S. Gopi, M. Javaheripi, P. Kauffmann, G. de Rosa, O. Saarikivi, A. Salim, S. Shah, H. S. Behl, X. Wang, S. Bubeck, R. Eldan, A. T. Kalai, Y. T. Lee, and Y. Li. Textbooks are all you need, 2023.
- [11] J. Hoffmann et al. An empirical analysis of compute-optimal large language model training. In A. H. Oh, A. Agarwal, D. Belgrave, and K. Cho, editors, *Advances in Neural Information Processing Systems*, 2022. URL <https://openreview.net/forum?id=iBBcRUOAPR>.
- [12] A. ImaniGooghari, P. Lin, A. H. Kargaran, S. Severini, M. Jalili Sabet, N. Kassner, C. Ma, H. Schmid, A. Martins, F. Yvon, and H. Schütze. Glot500: Scaling multilingual corpora and language models to 500 languages. In A. Rogers, J. Boyd-Graber, and N. Okazaki, editors, *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 1082–1117, Toronto, Canada, July 2023. Association for Computational Linguistics. doi: 10.18653/v1/2023.acl-long.61. URL <https://aclanthology.org/2023.acl-long.61>.
- [13] A. Q. Jiang, A. Sablayrolles, A. Mensch, C. Bamford, D. S. Chaplot, D. de Las Casas, F. Bressand, G. Lengyel, G. Lample, L. Saulnier, L. R. Lavaud, M. Lachaux, P. Stock, T. L. Scao, T. Lavril, T. Wang, T. Lacroix, and W. E. Sayed. Mistral 7b. *CoRR*, abs/2310.06825, 2023.
- [14] S. Lankford, H. Afli, and A. Way. Machine translation in the covid domain: an English-Irish case study for LoResMT 2021. In J. Ortega, A. K. Ojha, K. Kann, and C.-H. Liu, editors, *Proceedings of the 4th Workshop on Technologies for MT of Low Resource Languages (LoResMT2021)*, pages 144–150, Virtual, Aug. 2021. Association for Machine Translation in the Americas. URL <https://aclanthology.org/2021.mtsummit-loresmt.15>.
- [15] R. Luukkonen et al. FinGPT: Large generative models for a small language. In H. Bouamor, J. Pino, and K. Bali, editors, *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 2710–2726, Singapore, Dec. 2023. Association for Computational Linguistics. doi: 10.18653/v1/2023.emnlp-main.164. URL <https://aclanthology.org/2023.emnlp-main.164>.
- [16] P. Maini, S. Seto, H. Bai, D. Grangier, Y. Zhang, and N. Jaitly. Rephrasing the web: A recipe for compute and data-efficient language modeling, 2024.
- [17] S. Mille, E. Uí Dhonnchadha, L. Cassidy, B. Davis, S. Dasiopoulou, and A. Belz. Generating Irish text with a flexible plug-and-play architecture. In M. Surdeanu, E. Riloff, L. Chiticariu, D. Freitag, G. Hahn-Powell, C. T. Morrison, E. Noriega-Atala, R. Sharp, and M. Valenzuela-Escarcega, editors, *Proceedings of the 2nd Workshop on Pattern-based Approaches to NLP in the Age of Deep Learning*, pages 25–42, Singapore, Dec. 2023. Association for Computational Linguistics. doi: 10.18653/v1/2023.pandl-1.4. URL <https://aclanthology.org/2023.pandl-1.4>.
- [18] S. Mishra, D. Khashabi, C. Baral, and H. Hajishirzi. Cross-task generalization via natural language crowdsourcing instructions. In S. Muresan, P. Nakov, and A. Villavicencio, editors, *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 3470–3487, Dublin, Ireland, May 2022. Association for Computational Linguistics. doi: 10.18653/v1/2022.acl-long.244. URL <https://aclanthology.org/2022.acl-long.244>.
- [19] N. Muennighoff, A. M. Rush, B. Barak, T. L. Scao, N. Tazi, A. Piktus, S. Pyysalo, T. Wolf, and C. Raffel. Scaling data-constrained language models. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023. URL <https://openreview.net/forum?id=j5BuTrEj35>.
- [20] T. Nguyen, C. V. Nguyen, V. D. Lai, H. Man, N. T. Ngo, F. Dernoncourt, R. A. Rossi, and T. H. Nguyen. Culturax: A cleaned, enormous, and multilingual dataset for large language models in 167 languages, 2023.
- [21] L. Ouyang et al. Training language models to follow instructions with human feedback. In A. H. Oh, A. Agarwal, D. Belgrave, and K. Cho, editors, *Advances in Neural Information Processing Systems*, 2022. URL <https://openreview.net/forum?id=TG8KACxEOE>.
- [22] C. Raffel, N. M. Shazeer, A. Roberts, K. Lee, S. Narang, M. Matena, Y. Zhou, W. Li, and P. J. Liu. Exploring the limits of transfer learning with a unified text-to-text transformer. *J. Mach. Learn. Res.*, 21: 140:1–140:67, 2019. URL <https://api.semanticscholar.org/CorpusID:204838007>.
- [23] H. Touvron et al. Llama 2: Open foundation and fine-tuned chat models, 2023.
- [24] UNESCO, editor. *Atlas of the world’s languages in danger*. Memory of Peoples Series. UNESCO, 3 edition, Feb. 2010.
- [25] B. Workshop et al. Bloom: A 176b-parameter open-access multilingual language model, 2022.
- [26] H. Xu, Y. J. Kim, A. Sharaf, and H. H. Awadalla. A paradigm shift in machine translation: Boosting translation performance of large language models. In *The Twelfth International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=farT6XXntP>.
- [27] Z. X. Yong et al. BLOOM+1: Adding language support to BLOOM for zero-shot prompting. In A. Rogers, J. Boyd-Graber, and N. Okazaki, editors, *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 11682–11703, Toronto, Canada, July 2023. Association for Computational Linguistics. doi: 10.18653/v1/2023.acl-long.653. URL <https://aclanthology.org/2023.acl-long.653>.
- [28] L. Zheng, W.-L. Chiang, Y. Sheng, S. Zhuang, Z. Wu, Y. Zhuang, Z. Lin, Z. Li, D. Li, E. Xing, H. Zhang, J. E. Gonzalez, and I. Stoica. Judging LLM-as-a-judge with MT-bench and chatbot arena. In *Thirty-seventh Conference on Neural Information Processing Systems Datasets and Benchmarks Track*, 2023.