

Text cleaning with transformer language models for Hungarian

Gábor Madarász¹, András Holl², Noémi Ligeti-Nagy¹, Zijian Győző Yang¹,
Tamás Váradi¹

¹HUN-REN Hungarian Research Centre for Linguistics
surname.(middle name).forename@nytud.hun-ren.hu

²Library and Information Centre, Hungarian Academy of Sciences
holl.andras@konyvtar.mta.hu

Abstract. In language technology, clean data is fundamental for training high-quality models, yet large corpora often contain substantial noise due to OCR errors, missing diacritics, and various user-generated inconsistencies. This paper presents a comprehensive text cleaning pipeline tailored for Hungarian, leveraging transformer-based language models optimized for three key tasks: OCR error correction, diacritic restoration, and filtering grammatically incorrect sentences. We introduce huT5, a Hungarian adaptation of the mT5 model, which significantly reduces model parameters and resource demands while maintaining strong performance on Hungarian-specific text cleaning tasks. The huT5 models were fine-tuned on carefully constructed Hungarian corpora for each task and benchmarked against state-of-the-art methods, demonstrating competitive results, particularly in OCR error correction and diacritic restoration. Our pipeline offers an efficient, freely accessible solution to enhance data quality for Hungarian NLP applications, setting a new standard in resource-efficient, language-specific text cleaning.

Keywords: Text cleaning, Transformer Language Models, Hungarian NLP, OCR correction, Diacritic restoration, huT5 model

1 Introduction

Corpus cleaning is an ongoing challenge in language technology. Clean data is essential for training language models, yet with the large datasets needed for deep neural networks, a significant portion often requires cleaning. Text data can be messy for various reasons. A prominent research area is the cleaning of OCR-processed texts, which are often prone to errors. Encoding issues, missing diacritics, and grammatical mistakes due to user habits or conventions are also common issues. Cleaning these issues is complex, and no single solution fits all cases. While large language models like ChatGPT hold promise for handling many tasks, they are not perfect and often demand substantial resources, which can be impractical in some scenarios.

To tackle these challenges, we are developing targeted models and applications that can efficiently handle various text cleaning tasks with minimal resources while achieving high accuracy.

Our research prioritizes three main tasks for Hungarian: OCR cleaning, diacritic restoration, and filtering grammatically incorrect texts.

This paper presents a collaborative project between the HUN-REN Hungarian Research Centre for Linguistics (HUN-REN NYTK) and the Library and Information Centre of the Hungarian Academy of Sciences (MTA KIK), aimed at making the contents of the REAL Repository more accessible to researchers and easier to curate and enhance for the MTA KIK.

2 Related Work

The importance of improving texts processed through Optical Character Recognition (OCR) is emphasized by the creation of specialized competitions, such as the post-OCR error correction challenge instituted at the International Conference on Document Analysis and Recognition (ICDAR) since 2017¹. These competitions have catalyzed progress in OCR error rectification by utilizing cutting-edge methodologies. The method that secured victory in 2019, termed Context-based Character Correction (CCC, Rigaud et al., 2019), integrated a convolutional neural network with a BERT model for the purpose of error identification, coupled with a bidirectional Long Short-Term Memory (LSTM) encoder-decoder framework incorporating an attention mechanism for error rectification. Other scholars, including Nguyen et al. (2020), adapted CCC by incorporating Named Entity Recognition (NER) for the identification phase and substituting the LSTM-based correction model with a neural machine translation architecture known as OpenNMT. Schaefer and Neudecker (2020) also employed a dual-phase methodology, utilizing a bidirectional LSTM for error identification and a separate LSTM-based model for rectification.

Contemporary advancements have increasingly centered around transformer-based methodologies. Duong et al. (2021) addressed OCR correction without employing a specific detection module, instead opting to train a transformer model on a dataset encompassing both accurately transcribed text and OCR-induced erroneous text. They produced synthetic OCR errors through automated techniques, including random error insertion and parallel corpus training, which involved pairing pristine texts with their artificially flawed counterparts. In a similar vein, Mei et al. (2016) confronted scenarios involving multiple OCR correction models by devising a statistical approach to rank correction candidates predicated on attributes such as Levenshtein distance and lexicon validation.

Pethő et al. (2024) conducted an extensive study on types of OCR errors specific to Hungarian. Building on this, Laki et al. (2022) used neural language models to correct OCR errors, developing a two-part system with a Context-based Character Correction (Nguyen et al., 2020) detection module and an encoder-decoder model fine-tuned to perform the corrections. They also created a manually annotated Gold Standard training corpus for model training.

Numerous studies have explored the examination and cleaning of erroneous Hungarian texts using modern language technology tools. For instance, Dömötör

¹ <https://sites.google.com/view/icdar2017-postcorrectionocr>

and Yang (2018) conducted research on Hungarian corpora to identify the prevalence of non-standard errors, focusing primarily on spoken language and personal subcorpora. Their findings highlighted that the most frequent errors include the omission of punctuation, diacritics, and capitalization. This observation aligns with the fact that Hungarian researchers have undertaken separate studies to address each of these issues. Concerning language correctness, the HuLU benchmark kit includes the Hungarian Corpus of Linguistic Acceptability (HuCOLA) (Ligeti-Nagy et al., 2024), which specifically assesses the grammatical accuracy of Hungarian sentences.

In Hungarian punctuation research, Tündik et al. (2018) addressed punctuation correction using RNN networks, while Yang (2021) approached it with a transformer-based machine translation method.

Diacritic restoration in Hungarian has also been a longstanding area of research. Németh et al. (2000) presented a text-to-speech application that handled words without accents. Later, Ács and Halmai (2016) developed an n-gram-based statistical system that restores diacritics without relying on language-specific dictionaries. Several solutions use machine translation techniques to restore accents: Novák and Siklósi (2015) employed statistical machine translation, while Laki and Yang (2020a,b) trained neural machine translation models to solve this problem for Hungarian and other languages.

Despite the progress, most of these models are either not freely accessible or do not align well with our dataset, necessitating the development of custom models to meet our specific needs.

3 Corpora

In our task, we needed to clean a variety of different types of documents in a given library. The Library of the Hungarian Academy of Sciences was established in 1826 and has been serving the members of the Academy and the broader Hungarian research community ever since. The digital collections – in the form of an open access repository – were created in 2008. This repository – named REAL – has diverse holdings, mirroring the printed collection of the library. Its collection includes materials from multiple sources, such as manuscripts, books, and scientific papers, available in formats like printed copies and digital-born versions. This diversity of input channels has resulted in a rich, mixed collection – spanning scanned documents, born-digital files, publishers’ PDFs, accepted manuscripts, as well as various handwritten documents and images. For this project, we plan to use a modern text corpus comprising about 1 billion words. The REAL Repository contains over 250,000 documents, approximately half of which are suitable for our work.

3.1 Training corpora for the cleaning models

For the OCR task, we constructed a custom corpus to train our language model. To ensure accurate error detection, the corpus includes a balanced mix of erroneous and error-free text segment pairs, with a distribution of 66.4% to 33.6%,

respectively. For this, we used the corresponding error-free electronic versions of the OCR-processed texts. Table 1 presents the main characteristics of the training and test corpora. The average Character Error Rate (CER) across the entire training dataset is 12.35%, and the Word Error Rate (WER) is 11.74%, measured against the reference data (error-free sentences).

Dataset		Segments	Tokens	Avg Length
Training Set	Source	1,374,665	56,551,620	43.05
	Target	-	47,029,140	35.80
Test Set	Source	6,780	124,401	18.35
	Target	-	115,217	16.99

Table 1. Training and test corpora characteristics for the OCR cleaning task

For diacritic restoration, we used the same corpora as Laki and Yang (2020a) in their research on Hungarian. They chose the online available parallel corpus, Open Subtitles², which contains texts written in 62 languages, including Hungarian. It consists of movie subtitles with many shorter, informal sentences in them. Table 2 shows the main characteristics of the training and test corpora. Diacritic characters make up 6.59% of the total characters, and 35.75% of the words contain at least one diacritic.

Dataset	Segments	Tokens	Avg Length
Training Set	28,704,830	177,588,069	6.19
Test Set	3,000	18,635	6.21

Table 2. Training and test corpora characteristics for the diacritic restoration task

For filtering incorrect sentences, we used the HuCOLA corpus from the HuLU benchmark collection (Ligeti-Nagy et al., 2024). The corpus contains 9076 sentences labelled for their grammaticality. Here we used the training and the validation set, with 7276 and 900 sentences, respectively. Table 3 presents the main characteristics of the training and test corpora, where the distribution of incorrect and correct sentences is 21.6% and 78.4%, respectively.

Dataset	Segments	Tokens	Avg Length
Training Set	7,276	50,068	6.88
Test Set	900	6,145	6.75

Table 3. Characteristics of the HuCOLA training and test corpora

² <http://opus.nlpl.eu/OpenSubtitles-v2018.php>

4 Models and Experiments

4.1 huT5 models

Based on related work (Laki et al., 2022; Laki and Yang, 2020a), we identified the encoder-decoder architecture (Vaswani et al., 2017) as the most effective for our needs. However, no pretrained encoder-decoder language models specifically for Hungarian were available. Consequently, we developed a custom model by adapting the mT5 (Xue et al., 2021) for Hungarian. We followed guidance on single-language adaptation from a multilingual T5 model³, using methods to prune redundant embeddings and reduce parameter counts with minimal quality loss. Additionally, we optimized vocabulary size based on Abdaoui et al. (2020). This process resulted in two Hungarian-specific mT5 models (huT5) in both base and large versions, which will be made available on our Huggingface site⁴.

Once the huT5 models were created, we fine-tuned them for OCR cleaning, diacritic restoration, and incorrect sentence filtering.

4.2 Text cleaning pipeline

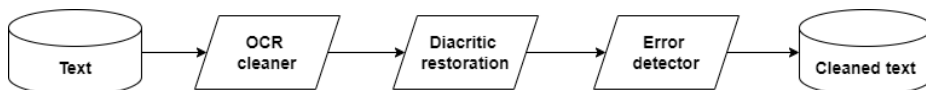


Fig. 1: Architecture of text cleaning pipeline

Using our fine-tuned, task-specific models, we can build a text cleaning pipeline to address our task. Figure 1 shows the architecture of our text cleaning pipeline. In our task, the primary errors originate from OCR; therefore, our first module is an OCR cleaner. The second major source of errors is missing or incorrect diacritics, so our second module is a diacritic restoration model. After these two modules complete the cleaning, if errors still remain, we use an erroneous sentence detector to mark or filter out the incorrect sentences.

In the next section, we will provide a detailed introduction to each of these modules.

4.3 Modules and models

For the OCR cleaning task, we initially tested the model from Laki et al. (2022), but it did not perform consistently across all text types in our dataset. Similarly,

³ <https://towardsdatascience.com/how-to-adapt-a-multilingual-t5-model-for-a-single-language-b9f94f3d9c90>

⁴ <https://huggingface.co/NYTK>

previous models for diacritic restoration (Laki and Yang, 2020a) were not freely available, prompting us to train custom models for these tasks.

Therefore, for these tasks, we fine-tuned custom, task-specific models based on our huT5 models. The large models were trained using two NVIDIA A100 GPUs (80GB each), while the base models were trained on four NVIDIA GeForce GTX 1080 GPUs (11GB each). We utilized the PyTorch Seq2SeqTrainer from the Transformers library⁵ for fine-tuning. For comparison, we also trained the original mT5 models on these tasks.

The training hyperparameters for the OCR cleaning models are as follows: learning rate = 5e-5; global batch size = 256; epoch = 1; sequence length = 256.

The training hyperparameters for the diacritic restoration models are as follows: learning rate = 5e-5; global batch size = 512; epoch = 1; sequence length = 256.

For both OCR cleaning and diacritic restoration tasks, we tailored the batch sizes to the available GPU and used gradient accumulation to achieve the global batch size. Due to the larger corpus size in diacritic restoration, we selected a larger global batch size for this task.

The training hyperparameters for the HuCOLA models are as follows: learning rate = 5e-6; global batch size = 32; epoch = 10; sequence length = 128. In the HuCOLA task, we trained for 10 epochs to ensure reliable comparisons, selecting the best-performing checkpoint. In the case of the HuCOLA experiment, we did not use gradient accumulation. For better comparison, the global batch size was kept the same as in the experiment conducted by Yang et al. (2023).

In the final text cleaning pipeline, the best-performing models constitute the OCR cleaning, diacritic restoration, and erroneous filtering modules.

To evaluate our huT5 models comprehensively, we also fine-tuned both huT5 and mT5 models on the Hungarian Corpus of Linguistic Acceptability (HuRTE) and Hungarian version of the Stanford Sentiment Treebank (HuRTE) benchmarks (Ligeti-Nagy et al., 2024).

5 Results

Our first objective was to evaluate the huT5 models. Table 4 provides a comparison between our huT5 and mT5 models. A key result is the significant reduction in both the number of parameters (base: ~42%; large: ~68%) and model size (base: ~42%; large: ~67%) following the conversion. For quality evaluation, we used the HuCOLA, HuSST, and HuRTE benchmarks, along with accuracy as the evaluation metric. The primary focus here is on the performance difference between the two model types rather than on outperforming existing Hungarian models.

The results indicate that in all cases, the huT5 models perform similarly or better than the mT5 models, despite having fewer parameters.

Consistent with findings from Yang (2022); Yang and Váradi (2021), the encoder-decoder architecture tends to be more suitable for sequence-to-sequence

⁵ <https://github.com/huggingface/transformers/tree/main/examples/pytorch>

tasks than for classification. This is reflected in our experiments with Hungarian benchmarks, where the models showed lower performance on classification tasks.

Model	Parameters	Size	HuCOLA	HuRTE	HuSST
huT5 base	244 million	977 MB	80.98	55.00	58.37
mT5 base	580 million	2.33 GB	80.98	52.66	57.94
huT5 large	820 million	3.28 GB	80.88	54.00	58.79
mT5 large	1.2 billion	4.92 GB	80.88	53.50	58.02

Table 4. Comparison of huT5 and mT5 on Hungarian benchmarks

Given that encoder-decoder models generally perform better on text generation tasks, we expect this model to be well-suited for our OCR cleaning and diacritic restoration tasks. In these experiments, we compared our models with Hungarian state-of-the-art (SOTA) models in this field, with benchmark results available in Laki et al. (2022) and Laki and Yang (2020a). The results can be seen in Table 5.

For evaluation, we used ROUGE-L See et al. (2017) (where higher scores indicate better quality $\uparrow$), Word Error Rate (WER, with lower values indicating better performance $\downarrow$), and Character Error Rate (CER, with lower values indicating better performance $\downarrow$) Morris et al. (2004) for OCR cleaning. For diacritic restoration, we assessed model performance using precision and recall metrics.

As you can see, our huT5 large model achieved the best performance in OCR cleaning, outperforming the SOTA model. In the diacritic restoration task, none of our models surpassed the SOTA models, but we achieved competitive results and will publish our models. Notably, our huT5 models, with fewer parameters, outperformed or achieved competitive performance compared to the mT5 models.

Model	OCR Cleaning			Diacritic Restoration	
	ROUGE-L $\uparrow$	WER $\downarrow$	CER $\downarrow$	Precision $\uparrow$	Recall $\uparrow$
SOTA model	93.44	17.38	8.72	99.38	99.28
huT5 base	95.12	11.10	7.05	97.64	97.75
mT5 base	95.26	10.61	6.57	97.57	97.49
huT5 large	95.66	10.01	6.46	97.95	98.18
mT5 large	95.28	11.56	9.33	98.61	98.61

Table 5. Comparison of OCR cleaning and diacritic restoration results for huT5 and mT5 models

The third module is an erroneous sentence detector. After the previous two modules have corrected the text, any remaining errors can, depending on the task, be marked or filtered out by an error detector. For this task, the HuCOLA benchmark is the most suitable. As shown in Table 4, our model achieved only

about 80% accuracy. In comparison, the models trained by Yang et al. (2023) performed better, with the HuBERT and PULI models achieving approximately 90% and 91% accuracy, respectively. Consequently, in our final pipeline, we use the fine-tuned PULI Bert-Large model instead of our fine-tuned huT5 or mT5.

We also performed error analysis on the OCR and diacritic restoration models, which achieved the best performance.

In Table 6, the error analysis of the diacritic restoration model is presented. The typical errors are similar to those reported by Laki and Yang (2020a), with only the ratios differing. In our case, we observed a higher proportion of real errors.

There are two main categories included in the correct output: equivalent forms and correct replaceable outputs. The equivalent form (e.g., *hova-hová* (where), *tied-tiéd* (yours)) refers to cases where both forms of the given word are usable, with no difference in meaning. The replaceable outputs refer to cases where the given words have different meanings, but both are correct either in the given context or without additional context. In the case of real errors, approximately 30% of the errors stem from proper nouns, as the model is unable to correctly restore names.

Error type	Ratio	Examples (reference (ref) - prediction (pred))
Correct output	37.5%	
Equivalent form		<i>hova - hívá</i> ('where'), <i>tied - tiéd</i> ('yours')
Replaceable output		ref: <i>Fogjátok meg!</i> ('Catch him/her!')
		pred: <i>Fogjatok meg!</i> ('Catch me!')
		ref: <i>Az InStyle magazinnal.</i> ('With InStyle magazine'.)
		pred: <i>Az InStyle magazinnál.</i> ('At InStyle magazine.')
Wrong reference		ref: <i>Két kijárat, egy elől</i> ('from'), <i>egy hátul.</i>
		pred: <i>Két kijárat, egy elől</i> ('in front'), <i>egy hátul.</i>
Real errors	62.5%	
Proper noun		<i>Livíró - Liuról, Ramával - Rámával</i>
Wrongly replaced		<i>még - meg, teli - téli</i>

Table 6. Error analysis of the diacritic restoration model

In Table 7, the error analysis of the OCR cleaning model is presented. Notably, 21% of the reported errors are actually false positives, where the reference text was incorrect, and the model provided the correct output. The remaining 79% of the errors represent real mistakes. These errors can be classified into three main categories: insertion, deletion, and replacement:

- Insertion: This occurs when the model adds an incorrect character to a word. Common cases involve the addition of extra punctuation marks or entirely incorrect letters. Additionally, the model may erroneously insert a space within a word, leading to segmentation errors. Finally, there are cases where the model inserts words that did not originally exist in the text.

- Deletion: In this category, the model removes a character that should have been retained. Similar to the insertion category, these deletions can involve punctuation marks, letters, spaces, or entire words. Deletion of spaces often results in the unintended merging of two words, while the deletion of whole words leads to a loss of information.
- Replacement: This category encompasses two primary types of errors: the substitution of an incorrect punctuation mark or letter. A notable subtype of replacement errors involves the mishandling of diacritic marks, where the model incorrectly replaces accented characters. Additionally, there are cases where the model replaces an entire word with a completely unrelated one, resulting in significant semantic deviations.

This analysis highlights specific patterns of error generation, providing insights for targeted model improvements.

Error type	Ratio	Examples (reference (ref) - prediction (pred))
Correct output	21.38%	<i>pcdig - pedig</i> <i>minder.kinek - mindenkinek</i> <i>színház- színház</i>
Real errors	78,62%	
Insertion	23,83%	<i>írva - írva:</i> (punctuation mark) <i>darab - dakrab</i> (letter) <i>mennybéli - menny béli</i> (space) <i>németalföldi - német 376 alföldi</i> (words)
Deletion	24,44%	<i>halastó, - halastó</i> (punctuation mark) <i>vesszővel - vesszvel</i> (letter) <i>már most - mármost</i> (space) <i>Hozott Isten! - []</i> (words)
Replacement	30,35%	<i>bort, - bort;</i> (punctuation mark) <i>előbb - elébb</i> (letter) <i>színes - színes</i> (diacritic) <i>Mi - Ali</i> (word)

Table 7. Error analysis of the OCR cleaning

6 Challenges and limitations

Adapting mT5 to huT5 presented several challenges, particularly in addressing the unique linguistic features of Hungarian, such as complex morphology and extensive diacritic use. These characteristics required substantial pre-processing to enhance model robustness, especially for tasks like diacritic restoration and grammatical error detection. Additionally, as Hungarian is relatively low-resource, the limited availability of high-quality annotated datasets constrained training, which may affect model performance on domain-specific or informal text corpora.

Another challenge was the computational demand required for model adaptation and fine-tuning. Despite efforts to reduce parameter count, the large huT5 version remains resource-intensive, potentially limiting its accessibility to groups without advanced hardware. Additionally, adapting huT5 to specific tasks required careful tuning to avoid overfitting due to dataset limitations.

Lastly, current evaluation metrics like WER, CER, and ROUGE-L offer a partial view of improvements in text quality and readability. Developing refined evaluation metrics tailored to Hungarian text cleaning could help better capture the model’s impact.

7 Conclusion

In this study, we addressed the essential task of text cleaning for Hungarian using custom transformer-based models. The results of our OCR cleaning, diacritic restoration, and incorrect sentence filtering tasks highlight the effectiveness and adaptability of our huT5 models, showing both improved performance and resource efficiency over existing models. By adapting the mT5 model specifically for Hungarian, we achieved substantial reductions in parameter counts and model size, with a remarkable 42% reduction for the base model and 68% for the large model. This optimization not only maintained but also, in many cases, improved the model’s performance on Hungarian benchmarks, particularly in sequence-to-sequence tasks.

Our results confirm that the huT5 models are well-suited for a range of text cleaning tasks. Compared to the original mT5, the huT5 models consistently achieved better scores on Hungarian OCR cleaning and diacritic restoration, as shown by lower WER and CER values and higher ROUGE-L scores. Additionally, the high precision and recall in diacritic restoration and the solid performance across all metrics make the huT5 models a strong candidate for state-of-the-art solutions in Hungarian text cleaning.

This work also provides a valuable, freely accessible alternative for Hungarian-language tasks, meeting gaps left by existing models. The encoder-decoder architecture used in our approach effectively addresses both sequence-to-sequence and error detection needs, presenting a refined tool for improving data quality in Hungarian NLP applications.

Future work may extend these results by testing the huT5 models’ adaptability to additional languages or dialects. Nonetheless, this study establishes a new standard for Hungarian text cleaning, showing that transformer-based approaches, when tailored to the specific language requirements, can achieve both high accuracy and efficiency.

Acknowledgment

The research was supported by the MTA “Tudomány a Magyar Nyelvért Nemzeti Program” (Science for the Hungarian Language National Programme).

References

- Abdaoui, A., Pradel, C., Sigel, G.: Load what you need: Smaller versions of multilingual BERT. In: Moosavi, N.S., Fan, A., Shwartz, V., Glavaš, G., Joty, S., Wang, A., Wolf, T. (eds.) *Proceedings of SustaiNLP: Workshop on Simple and Efficient Natural Language Processing*. pp. 119–123. Association for Computational Linguistics, Online (Nov 2020), <https://aclanthology.org/2020.sustainlp-1.16>
- Duong, Q., Härmäläinen, M., Hengchen, S.: An unsupervised method for OCR post-correction and spelling normalisation for Finnish. In: *Proceedings of the 23rd Nordic Conference on Computational Linguistics (NoDaLiDa)*. pp. 240–248. Linköping University Electronic Press, Sweden, Reykjavik, Iceland (Online) (May 31–2 Jun 2021)
- Dömötör, A., Yang, Z.Gy.: Így írtok ti: nem sztenderd szövegek hibatípusainak detektálása gépi tanulásos módszerrel [This is how you write: detecting error types in non-standard texts using machine learning]. XIV. Magyar Számítógépes Nyelvészeti Konferencia pp. 305–316 (2018)
- Laki, L.J., Kőrös, A., Ligeti-Nagy, N., Nyéki, B., Vadász, N., Yang, Z.Gy., Váradi, T.: OCR-hibák javítása neurális technológiák segítségével [Correcting OCR errors using neural technologies]. In: XVIII. Magyar Számítógépes Nyelvészeti Konferencia (MSZNY 2022). pp. 417–430. Szegedi Tudományegyetem, Informatikai Intézet, Szeged, Hungary (2022)
- Laki, L.J., Yang, Z.Gy.: Automatic Diacritic Restoration With Transformer Model Based Neural Machine Translation for East-Central European Languages. In: *Proceedings of the 11th International Conference on Applied Informatics (ICAI 2020)*. pp. 190–202. Eger, Hungary (2020a), <http://ceur-ws.org/Vol-2650/>
- Laki, L.J., Yang, Z.Gy.: Automatikus ékezetvisszaállítás transzformer modellen alapuló neurális gépi fordítással [Automatic accent restoration using neural machine translation based on transformer model]. XVI. Magyar Számítógépes Nyelvészeti Konferencia pp. 181–190 (2020b)
- Ligeti-Nagy, N., Ferenczi, G., Héja, E., Laki, L.J., Vadász, N., Yang, Z.G., Váradi, T.: HuLU: Hungarian language understanding benchmark kit. In: Calzolari, N., Kan, M.Y., Hoste, V., Lenci, A., Sakti, S., Xue, N. (eds.) *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*. pp. 8360–8371. ELRA and ICCL, Torino, Italia (May 2024), <https://aclanthology.org/2024.lrec-main.733>
- Mei, J., Islam, A., Wu, Y., Moh’d, A., Milios, E.E.: Statistical learning for OCR text correction. CoRR abs/1611.06950 (2016), <http://arxiv.org/abs/1611.06950>
- Morris, A., Maier, V., Green, P.: From WER and RIL to MER and WIL: improved evaluation measures for connected speech recognition. pp. 2765–2768 (10 2004)
- Németh, G., Zainkó, Cs., Fekete, L., Olaszy, G., Endrédi, G., Olaszi, P., Kiss, G., Kis, P.: The design, implementation, and operation of a hungarian e-mail

- reader. *International Journal of Speech Technology* 3(3), 217–236 (Dec 2000), <https://doi.org/10.1023/A:1026567216832>
- Nguyen, T.T.H., Jatowt, A., Nguyen, N.V., Coustaty, M., Doucet, A.: Neural Machine Translation with BERT for Post-OCR Error Detection and Correction. In: *Proceedings of the ACM/IEEE Joint Conference on Digital Libraries in 2020*. p. 333–336. JCDL '20, Association for Computing Machinery, New York, NY, USA (2020)
- Novák, A., Siklósi, B.: Automatic diacritics restoration for Hungarian. In: *Proceedings of the 2015 Conference on Empirical Methods in Natural Language Processing*. pp. 2286–2291. Association for Computational Linguistics, Lisbon, Portugal (Sep 2015), <https://www.aclweb.org/anthology/D15-1275>
- Pethő, G., Sass, B., Simon, L., Lipp, V.: OCR-hibák kvantitatív elemzése több szövegváltozat összehasonlításával [Quantitative analysis of OCR errors through the comparison of multiple text versions]. In: *XX. Magyar Számítógépes Nyelvészeti Konferencia (MSZNY 2024)*. pp. 17–29. Szegedi Tudományegyetem, Informatikai Intézet, Szeged, Hungary (2024)
- Rigaud, C., Doucet, A., Coustaty, M., Moreux, J.P.: Icdar 2019 competition on post-ocr text correction. In: *2019 International Conference on Document Analysis and Recognition (ICDAR)*. pp. 1588–1593 (2019)
- Schaefer, R., Neudecker, C.: A two-step approach for automatic OCR post-correction. In: *Proceedings of the The 4th Joint SIGHUM Workshop on Computational Linguistics for Cultural Heritage, Social Sciences, Humanities and Literature*. pp. 52–57. International Committee on Computational Linguistics, Online (Dec 2020)
- See, A., Liu, P.J., Manning, C.D.: Get to the point: Summarization with pointer-generator networks. In: Barzilay, R., Kan, M.Y. (eds.) *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*. pp. 1073–1083. Association for Computational Linguistics, Vancouver, Canada (Jul 2017), <https://aclanthology.org/P17-1099>
- Tündik, M.A., Tarján, B., Szaszák, G.: Televíziós feliratok írásjeleinek visszaállítása rekurrens neurális hálózatokkal [Restoration of punctuation marks in television subtitles using Recurrent Neural Networks]. *XIV. Magyar Számítógépes Nyelvészeti Konferencia* pp. 183–195 (2018)
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A.N., Kaiser, L.u., Polosukhin, I.: Attention is All you Need. In: Guyon, I., Luxburg, U.V., Bengio, S., Wallach, H., Fergus, R., Vishwanathan, S., Garnett, R. (eds.) *Advances in Neural Information Processing Systems* 30, pp. 5998–6008. Curran Associates, Inc. (2017), <http://papers.nips.cc/paper/7181-attention-is-all-you-need.pdf>
- Xue, L., Constant, N., Roberts, A., Kale, M., Al-Rfou, R., Siddhant, A., Barua, A., Raffel, C.: mT5: A massively multilingual pre-trained text-to-text transformer. In: Toutanova, K., Rumshisky, A., Zettlemoyer, L., Hakkani-Tur, D., Beltagy, I., Bethard, S., Cotterell, R., Chakraborty, T., Zhou, Y. (eds.) *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*.

- pp. 483–498. Association for Computational Linguistics, Online (Jun 2021), <https://aclanthology.org/2021.naacl-main.41>
- Yang, Z.Gy.: Automatikus írásjelek visszaállítása és Nagybetűsítés statikus korpuszon, transzformer modellen alapuló neurális gépi fordítással [Automatic punctuation restoration and capitalization on a static corpus using transformer-based neural machine translation]. pp. 225–232. Szegedi Tudományegyetem, Informatikai Tanszékcsoport, Szeged (2021)
- Yang, Z.Gy.: BARTerezzünk! - Messze, messze, messze a világtól, - BART kísérleti modellek magyar nyelvre [Let's BART! - Far, far, far away from the world, - BART experimental models for Hungarian]. In: XVIII. Magyar Számítógépes Nyelvészeti Konferencia. pp. 15–29. Szegedi Tudományegyetem, Informatikai Intézet, Szeged, Magyarország (2022)
- Yang, Z.Gy., Dodé, R., Ferenczi, G., Héja, E., Jelencsik-Mátyus, K., Kőrös, A., Laki, L.J., Ligeti-Nagy, N., Vadász, N., Váradi, T.: Jönnek a nagyok! BERT-Large, GPT-2 és GPT-3 nyelvmodellek magyar nyelvre [The giants are coming! BERT-Large, GPT-2, and GPT-3 language models for Hungarian]. In: XIX. Magyar Számítógépes Nyelvészeti Konferencia (MSZNY 2023). pp. 247–262. Szegedi Tudományegyetem, Informatikai Intézet, Szeged, Hungary (2023)
- Yang, Z.Gy., Váradi, T.: Training language models with low resources: RoBERTa, BART and ELECTRA experimental models for Hungarian. In: Proceedings of 12th IEEE International Conference on Cognitive Infocommunications (CogInfoCom 2021). pp. 279–285. IEEE, Online (2021)
- Ács, J., Halmi, J.: Hunaccent: Small Footprint Diacritic Restoration for Social Media. In: Proceedings of the Tenth International Conference on Language Resources and Evaluation (LREC 2016). pp. 3526–3529. Portoroz, Slovenia (2016)