

ADAPTING LARGE LANGUAGE MODELS FOR BULGARIAN CLINICAL DIETETICS: A MULTI-STAGE TRAINING APPROACH

Simeon Monov, Detelinka Trifonova, Nikolay Pavlov

Abstract. *The application of Large Language Models (LLMs) in specialized medical domains remains challenging, especially for low-resource languages such as Bulgarian. Clinical dietetics requires the interpretation of medical terminology, laboratory data, patient-specific restrictions, and therapeutic nutrition principles. This paper presents a multi-stage approach for adapting a general-purpose LLM to Bulgarian clinical dietetics. The proposed methodology includes the development of three specialized datasets: an educational corpus, a question-answering dataset, and structured clinical dietetic cases represented in JSON format. These data were used in a training pipeline consisting of Domain-Adaptive Pre-Training (DAPT), Supervised Fine-Tuning (SFT) for Q&A tasks, and SFT for structured diet generation. The model was trained using parameter-efficient methods and 4-bit quantization on a single NVIDIA Tesla V100 GPU with 32 GB VRAM. Final evaluation was performed on 114 unseen examples using structural and content-based metrics. The results show promising potential for a proof-of-concept assistant supporting Bulgarian clinical dietetic practice.*

Key Words: *Large Language Models, Clinical Dietetics, Bulgarian Language, Low-Resource Languages, DAPT, SFT, JSON, Medical AI.*

1. Introduction

Large Language Models have significantly changed natural language processing by enabling coherent text generation, question answering, summarization, and decision-oriented workflows. However, their application in specialized medical domains remains limited. This limitation becomes more visible when the target language is low-resourced, where domain-specific corpora, benchmarks, and adapted models are scarce.

Clinical dietetics is a demanding field because it combines medical terminology, nutritional science, laboratory interpretation, patient-specific constraints, and practical dietary recommendations. A dietetic recommendation is

more than a free-text answer; it often requires the interpretation of diagnoses, comorbidities, allergies, dietary restrictions, cultural preferences, and therapeutic goals. A general-purpose model may produce fluent output, but fluency alone does not ensure domain adequacy or structural consistency.

The motivation of this work is the lack of specialized Bulgarian-language datasets and models for clinical dietetics. While English-language medical NLP resources are more widely available, Bulgarian medical and dietetic corpora remain fragmented and difficult to use for model adaptation. This creates a gap between modern LLM capabilities and the practical needs of Bulgarian medical education and dietetic practice.

To address this gap, this paper presents a multi-stage approach for adapting a general-purpose LLM to Bulgarian clinical dietetics. The central idea is that a model should first be exposed to domain-specific educational material through Domain-Adaptive Pre-Training and then further specialized through supervised training on practical tasks. The main contributions are the creation of specialized datasets, the implementation of a three-stage adaptation pipeline, the use of structured JSON outputs for diet generation, and an expert-oriented evaluation of the model as a dietetic assistant.

2. Related Work and Motivation

Recent research shows the increasing potential of LLMs in medicine, clinical nutrition, and health-related decision support. Medical domain adaptation is important because general-purpose LLMs may lack sufficient terminology control, factual grounding, and clinical reliability. Chen et al. [1] present HuatuoGPT-II as an example of medical LLM adaptation. Belkhouribchia and Pen [2] discuss LLM applications, limitations, and future prospects in clinical nutrition. Niszczoła et al. [3] emphasize the need for structured prompts, model validation, and expert oversight in real-world nutrition assessment.

Most available medical language resources are oriented toward English. For Bulgarian, the availability of specialized medical corpora is considerably lower. In clinical dietetics, the limitation is even more specific because the domain requires medical terminology, local food culture, therapeutic diet principles, and the ability to produce structured recommendations. Therefore, adaptation should combine domain knowledge acquisition with task-specific supervision.

Domain-Adaptive Pre-Training has been proposed as an effective strategy for improving model performance when the target domain differs from the original training data [4]. Parameter-efficient fine-tuning methods such as LoRA [5] and QLoRA [6] are also relevant because they make it possible to adapt large models with limited computational resources. The proposed approach combines these ideas in a practical pipeline for Bulgarian clinical dietetics.

3. Dataset Development

The dataset development process was a central part of the project. The goal was not only to train a model, but also to construct a domain-specific data foundation for Bulgarian clinical dietetics. Three dataset types were prepared: an educational text corpus, a question-answering dataset, and a structured clinical diet generation dataset.

The first dataset consists of educational materials from a three-year training program in dietetic nutrition. The materials include textbooks, lecture notes, methodological resources, and specialized educational content. The documents were collected, filtered, cleaned, and converted into Markdown format. After preprocessing, the educational corpus used for DAPT was represented as 2,825 training examples. Markdown was selected because it preserves structure while remaining readable both by humans and machines. Headings, lists, tables, and emphasized terms are represented in a lightweight and predictable way, which is useful for language modeling.

The conversion process included cleaning OCR errors, removing duplicate content, standardizing terminology, and separating images from text. Figures and graphical materials were stored separately for possible future multimodal extensions. This step is important because clinical dietetics often involve tabular and visual information, including laboratory values, menus, and therapeutic diet schemes.

The second dataset was developed for question-answering tasks. Based on the educational corpus, a large set of open-ended Q&A examples was generated. The developed Q&A dataset contains 9,326 open-answer pairs used as an empirical basis for further model training. These examples teach the model to respond to practical questions in clinical dietetics and transform it from a passive language model into an interactive assistant.

The third dataset was designed for structured diet generation. It contains 5,343 JSON-formatted examples covering 11 clinical conditions, including hypertension, chronic kidney disease, liver disease, cardiovascular disease, gastritis, irritable bowel syndrome, chronic pancreatitis, hepatic steatosis, type 2 diabetes, iron-deficiency anemia, dyslipidemia, and hyperuricemia. The JSON schema represents patient characteristics, diagnoses, laboratory data, restrictions, dietary goals, and recommended meal structures. This reduces ambiguity and allows the generated output to be checked more systematically.

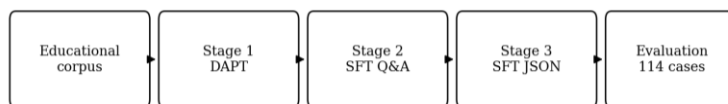
4. Model Selection and Training Pipeline

The selected base model was a Mistral-based 24B model in a quantized 4-bit version. The model was chosen after considering alternatives such as Mistral 7B, LLaMA-based models, and MedGemma. The larger 24B model offered a

better balance between capacity, stability, and domain adaptation potential, while 4-bit quantization made the training feasible on limited hardware.

The technical implementation used unsloth/Mistral-Small-3.2-24B-Instruct-2506-bnb-4bit with 4-bit quantization [7], based on the Mistral Small model family [8]. In the implemented training setup, the model was trained on an NVIDIA Tesla V100 GPU with 32 GB VRAM. The optimized quantized version made it possible to work with a 24B-parameter model within the available hardware constraints.

The training pipeline consisted of three stages: Domain-Adaptive Pre-Training on the educational corpus, Supervised Fine-Tuning for Q&A tasks, and Supervised Fine-Tuning for structured diet generation. This staged design follows the assumption that domain knowledge and task behavior should be learned progressively. DAPT exposes the model to the language of the domain; Q&A fine-tuning teaches it to answer domain-specific questions, and JSON fine-tuning teaches it to produce structured clinical dietetic recommendations.



Q&A dataset: 9,326 pairs; structured dataset: 5,343 JSON cases

Figure 1. Multi-stage training pipeline for Bulgarian clinical dietetics

4.1. Stage 1: Domain-Adaptive Pre-Training

The first stage used Domain-Adaptive Pre-Training to specialize the model in the language of Bulgarian clinical dietetics. This step continued the training of the base model on the curated educational corpus. The objective was causal language modeling, where the model learns to predict the next token based on the previous context. This objective is suitable for generative LLMs and supports later tasks such as answering questions and generating diet plans.

During DAPT, the model was exposed to specialized terms, dietetic principles, disease-related dietary restrictions, and educational explanations. The goal was not simple memorization, but adaptation of internal representations toward the terminology and semantic relations of the domain. The DAPT stage used parameter-efficient training with 1 GPU, 2,825 examples, 7 epochs, 623 total steps, a per-device batch size of 8, gradient accumulation of 4, a total batch size of 32, and 202,899,456 trainable parameters, corresponding to approximately 0.84% of the 24B model. The reported training duration was 14:40:20.

4.2. Stage 2: Supervised Fine-Tuning for Q&A

After DAPT, the model was fine-tuned on the Q&A dataset. The goal was to teach the model to use the acquired domain knowledge in an interactive

question-answering format. This is important because practical use cases in education and dietetic support often begin with natural language questions.

The examples were structured as prompts and responses, allowing the model to learn the relationship between user questions and clinically relevant answers. This stage focused on response relevance, terminology consistency, and domain-specific explanation quality. In practical terms, it helps the model answer questions about dietetic principles, disease-specific restrictions, and nutrition-related clinical reasoning. The Q&A SFT stage was performed on 9,326 open-answer pairs using the same quantized base model and parameter-efficient fine-tuning framework. It focused on conversational response quality rather than strict structural validation and was completed in approximately 3 hours on the same NVIDIA Tesla V100 GPU with 32 GB VRAM.

4.3. Stage 3: Supervised Fine-Tuning for Structured Diet Generation

The third stage was the most task-specific part of the methodology. It used structured JSON examples to train the model to generate dietetic recommendations in a machine-readable format. Unlike ordinary text generation, JSON generation requires both content accuracy and structural correctness. The model must produce valid syntax, preserve expected keys, and fill the fields with clinically meaningful information.

The structured diet generation task reflects the complexity of real dietetic decision-making. A diet plan must consider diagnosis, laboratory data, age, sex, comorbidities, allergies, restrictions, and dietary goals. Therefore, the model was trained not only to produce recommendations, but also to organize them into a format that can be inspected, validated, and reused in future software systems.

This stage used custom evaluation metrics such as `eval_json_valid_rate` and `eval_json_key_recall`, because standard loss alone is insufficient for evaluating structured generation. A response may have low loss but still fail as a JSON object; therefore, structural metrics are essential. The structured JSON SFT stage was performed on 5,343 JSON-formatted examples using the same quantized base model and parameter-efficient fine-tuning framework. In addition to content generation, this stage required the model to follow the predefined JSON structure. The training required approximately 21 hours on the same NVIDIA Tesla V100 GPU with 32 GB VRAM.

5. Evaluation

After the multi-stage training process, the final evaluation was performed on a separate hold-out test set consisting of 114 unseen clinical cases. The hold-

out test set was excluded from all training stages and used only for final evaluation.

Due to copyright, educational licensing, and research integrity constraints, the dataset is not publicly released. However, the evaluation protocol, including dataset structure, preprocessing steps, and metric computation methodology, is reproducible upon reasonable request.

The maximum number of new tokens for generation was set to 8192, allowing the model to generate complete structured responses for complex cases. Evaluation included both content-based and structure-based metrics to ensure clinical relevance and syntactic validity of the generated JSON outputs.

Table 1. Final evaluation metrics on the hold-out test set

Metric	Value	Interpretation
eval_loss	0.0888	Cross-entropy loss on the test set
eval_json_valid_rate	77.2%	Generated outputs that are valid JSON
eval_json_key_recall	50.4%	Valid JSON outputs containing correct top-level keys
eval_f1	50.0%	Token-overlap F1 score
eval_rougeL_f1	44.9%	ROUGE-L F1 based on longest common subsequence
eval_bleu	20.4%	N-gram precision-based generation metric

The most important result is the 77.2% JSON validity rate, which indicates that the model learned to produce syntactically valid structured outputs in most cases. This is a strong result for a complex generation task because the output is not merely a short answer, but a multi-field structured object.

The eval_json_key_recall value of 50.4% shows that the model partially learned the expected schema and that future work should focus on improving schema adherence. The content-based metrics also indicate useful performance. The eval_f1 value of 50.0% and eval_rougeL_f1 value of 44.9% suggest substantial overlap with reference outputs, while the lower eval_bleu value of 20.4% is expected for open-ended clinical generation tasks where multiple correct formulations may exist.

A qualitative evaluation was also conducted. The model demonstrated the ability to generate clinically relevant recommendations and to follow the general structure of dietetic reasoning. However, the analysis identified limitations, including a tendency toward repetitive menus and the need for further improvement in complex metabolic conditions. The model should therefore be interpreted as a proof-of-concept assistant, not as a replacement for a medical or dietetic specialist.

6. Discussion

The results demonstrate that a multi-stage adaptation strategy can transform a general-purpose model into a more domain-aware assistant for Bulgarian clinical dietetics. The most important aspect of the proposed pipeline is the separation of knowledge acquisition and task specialization. DAPT adapts the model to the domain language, Q&A SFT adapts it to explanatory interaction, and JSON SFT adapts it to structured clinical output.

This design is suitable for low-resource languages. In high-resource settings, large specialized datasets may already exist. In Bulgarian clinical dietetics, however, the absence of public corpora requires the creation of new datasets. The educational corpus provides foundational domain knowledge, while the Q&A and JSON datasets provide task-specific supervision.

The use of JSON is important from a methodological perspective. Many LLM applications in medicine rely on free-text responses, which are difficult to evaluate and integrate into software systems. Structured output makes the task more constrained and measurable. At the same time, structural validity alone is not sufficient. A valid JSON object can still contain incomplete, repetitive, or clinically weak recommendations, so automatic metrics should be combined with expert review.

The project also demonstrates that resource-efficient training is feasible. By using quantization and parameter-efficient fine-tuning, the model was trained on a single NVIDIA Tesla V100 GPU with 32 GB VRAM. The main limitations of the current model are related to dataset size, clinical diversity, output variability, and the risk of hallucination. Future systems should include retrieval-augmented generation, rule-based validation, and expert-in-the-loop review.

7. Conclusion

This paper presented a multi-stage approach for adapting a Large Language Model to Bulgarian clinical dietetics. The proposed methodology combines specialized dataset development with domain-adaptive and supervised training stages. The educational corpus was used for DAPT, the Q&A dataset for conversational fine-tuning, and the structured JSON dataset for practical diet generation.

The final evaluation showed that the adapted model can generate valid JSON outputs in 77.2% of cases and achieved promising content-based results, including 50.0% F1 and 44.9% ROUGE-L F1 on the test set. These results indicate that the proposed strategy is suitable for developing a proof-of-concept assistant for clinical dietetic practice in Bulgarian.

The main contribution of the work is not only the trained model, but also the methodology and datasets created for a low-resource medical domain. Future work will focus on expanding the datasets, improving schema consistency,

increasing menu diversity, integrating retrieval-augmented generation, and strengthening expert validation.

Acknowledgments

This study is supported by the project SP25-FMI-008 “Research and implementation of modern language models and artificial neural networks for automated processing, forecasting, and structuring of data and texts in specific application domains” at the Paisii Hilendarski University of Plovdiv. The authors also acknowledge the importance of expert support in the preparation and validation of the clinical dietetic materials used for dataset development.

References

- [1] J. Chen, X. Wang, K. Ji, et al., HuatuoGPT-II: One-stage training for medical adaption of LLMs, *arXiv*, 2023, ISSN: 2331-8422, DOI: <https://doi.org/10.48550/arxiv.2311.09774>
- [2] A. Belkhouribchia, C. Pen, Large language models in clinical nutrition: An overview of its applications, capabilities, limitations, and potential future prospects, *Frontiers in Nutrition*, 2025, vol. 12, ISSN: 2296-861X, DOI: <https://doi.org/10.3389/fnut.2025.1635682>
- [3] P. Niszczoła, I. Rybicka, M. Kaczmarek, Large language models for real-world nutrition assessment: Structured prompts, multi-model validation and expert oversight, *Nutrients*, 2026, vol. 18, no. 1, pp. 23–41, ISSN: 2072-6643, DOI: <https://doi.org/10.3390/nu18010023>
- [4] S. Gururangan, A. Marasović, S. Swayamdipta, K. Lo, I. Beltagy, D. Downey, N. Smith, Don’t Stop Pretraining: Adapt Language Models to Domains and Tasks, *Proc. of the 58th Annual Meeting of the Association for Computational Linguistics (ACL)*, 2020, pp. 8342–8360, DOI: 10.18653/v1/2020.acl-main.740
- [5] E. Hu, Y. Shen, P. Wallis, Z. Allen-Zhu, Y. Li, S. Wang, L. Wang, W. Chen, LoRA: Low-Rank Adaptation of Large Language Models, *International Conference on Learning Representations (ICLR)*, 2022, DOI: <https://doi.org/10.48550/arXiv.2106.09685>
- [6] T. Dettmers, A. Pagnoni, A. Holtzman, L. Zettlemoyer, QLoRA: Efficient Finetuning of Quantized LLMs, *Advances in Neural Information Processing Systems*, 2023, vol. 36, DOI: <https://doi.org/10.48550/arXiv.2305.14314>
- [7] Unsloth, “Mistral-Small-3.2-24B-Instruct-2506-bnb-4bit Model Card”, Hugging Face, available: <https://huggingface.co/unsloth/Mistral-Small-3.2-24B-Instruct-2506-bnb-4bit>
- [8] Mistral AI, “Mistral Small Model Documentation and Technical Description”, available: <https://docs.mistral.ai/>

Simeon Monov^{1, *}, Detelinka Trifonova¹, Nikolay Pavlov¹

¹ Paisii Hilendarski University of Plovdiv,

Faculty of Mathematics and Informatics,

236 Bulgaria Blvd., 4027 Plovdiv, Bulgaria

Corresponding author: smonov@uni-plovdiv.bg

