

Bosch@AI_Team at MMT 2025: Medical Machine Translation by Bidirectional Training with Small Language Models

Phan Minh Toan¹, Nguyen Xuan Phi¹, Nguyen Van Tai¹, Trang Minh Quang¹, Dang Van Thin²

¹Bosch Global Software Technologies, Ho Chi Minh City, Vietnam

²University of Information Technology-VNUHCM, Ho Chi Minh City, Vietnam

{quang.tranminh2,tai.nguyenvan,toan.phanminh}@vn.bosch.com

s3757281@rmit.edu.vn

external.Phi.NguyenXuan@bcn.bosch.com

thindv@uit.edu.vn

Abstract

Machine Translation (MT) for the medical domain poses unique challenges, as mistranslations of technical terminology can directly impact patient safety and hinder the accessibility of clinical knowledge. At the VLSP 2025 shared task on English–Vietnamese medical translation, we investigate whether small pre-trained language models can provide accurate and efficient domain-specific translations under realistic resource constraints. Our approach combines three key techniques: (1) prompt engineering with both direct and instruction templates to guide medical terminology usage, (2) bidirectional training to reinforce bilingual consistency, and (3) full-parameter fine-tuning of compact Qwen models (0.5–3B parameters). Experiments demonstrate that our Qwen3 1.7B model achieves substantial improvements over LoRA-based variants, particularly in the En–Vi direction. In the final private leaderboard, our system ranked second overall among participating teams, demonstrating that carefully adapted small language models, equipped with domain-specific prompting and training strategies that can surpass larger baselines in medical MT tasks.

1 Introduction

Machine Translation (MT) has advanced from rule-based systems to statistical models and now to neural architectures (Castilho and Knowles, 2025). Despite impressive gains in fluency and general-domain accuracy, the medical domain remains a critical unsolved problem. Errors in translating technical terminology can directly affect patient safety and hinder the dissemination of medical knowledge. Current MT systems, which are typically trained on broad web-scale corpora, often fail to capture specialized terminology and stylistic conventions, leading to mistranslations that reduce their reliability in clinical settings (Rios et al., 2022).

The problem is compounded by practical constraints. Large language models demand substantial computational resources, making them difficult to deploy in hospitals, universities, or NGOs with limited infrastructure. Yet, these are precisely the contexts where accurate translation is most needed. What is required, therefore, is a formulation of medical MT that emphasizes terminology fidelity, domain robustness, and efficiency of deployment. Prior work on lightweight and adaptable workflows, such as dictionary-assisted MT, demonstrates that resource-conscious solutions can still meet user needs in healthcare environments (Merx et al., 2025). This points toward the need for compact, specialized models that can operate under realistic constraints while delivering trustworthy medical translations

For Vietnamese, medical MT is both a linguistic and societal necessity. Healthcare professionals, students, and researchers depend heavily on English-language resources such as textbooks, clinical guidelines, and research articles, yet there is no consistent mechanism to make this knowledge reliably accessible in Vietnamese (Vo et al., 2024). This creates a persistent barrier to medical education, clinical decision-making, and research dissemination. Existing Vietnamese–English corpora, such as PhoMT and MTet, support general MT research but lack sufficient coverage of specialized medical texts (Doan et al., 2021; Ngo et al., 2022). As a result, current systems often fail to capture domain-specific terminology and style, leading to mistranslations that reduce trust and usability in clinical contexts. Moreover, the computational cost of training and deploying very large LLMs makes them impractical in many Vietnamese healthcare and academic settings.

This paper addresses existing limitations by focusing on improving Vietnamese medical machine translation through a methodology that leverages domain-specific pre-training and fine-tuning tech-

niques to enhance the performance of neural machine translation models. Specifically, we develop and evaluate several state-of-the-art neural machine translation models, including a fine-tuned version of a powerful pre-trained model tailored for the medical domain.

2 Related work

In this section, we focus on prior studies related to Vietnamese machine translation over the past three years, considering both general domains and the specialized context of the medical domain.

Machine Translation for Low-Resource Languages Low-resource Neural Machine Translation typically leverages cross-lingual transfer and synthetic data augmentation. Multilingual fine-tuning of pretrained models with iterative back-translation remains highly effective; for instance, EcXTra (Li et al., 2022) achieves strong performance on unseen pairs via this two-stage process. Recent work also explores hybrid architectures, such as a BART-based system for Bahnaric–Vietnamese (Nguyen et al., 2025), and improved back-translation through enhanced monolingual text quality (Pham et al., 2023). Overall, evidence suggests that careful data curation and domain-specific fine-tuning often matter more than increasing model size in low-resource MT.

Machine Translation in the Medical Domain

In high-risk domains such as medicine, accurate terminology and reliable evaluation are critical. Domain adaptation through fine-tuning multilingual pretrained models (e.g., mBART) has been shown to reduce terminology errors and improve adequacy, with MQM-based evaluation frameworks highlighting the limitations of automatic metrics (Gaona et al., 2023). Moreover, recent work on clinical Machine Translation shows that smaller multilingual pre-train language models, when carefully fine-tuned, can outperform very large models, reinforcing the importance of domain-specific adaptation over model scale (Han et al., 2024).

Vietnamese Medical Machine Translation Recent resources include MedEV (Vo et al., 2024), a manually validated corpus of $\sim 360,000$ En–Vi sentence pairs, and large synthetic datasets derived from translated PubMed abstracts used to pretrain ViPubmedT5, together with the human-refined ViMedNLI; these resources, plus self-

training and back-translation, substantially improve domain adaptation and BLEU for En–Vi biomedical tasks (Phan et al., 2022). MTet, a large multi-domain English–Vietnamese corpus containing biomedical data, was introduced together with En–ViT5, a bilingual pretrained model that achieved state-of-the-art performance (Ngo et al., 2022). Another study investigated clinical MT with multilingual pretrained models and demonstrated that domain-adaptive fine-tuning is essential for handling medical terminology (Han et al., 2024). Together, these works emphasize the need for high-quality resources and domain adaptation in Vietnamese medical machine translation.

3 Methodology

3.1 Prompt Engineering

We treat machine translation as an instruction following task by carefully designing natural-language prompts. Prompt engineering is known to be critical for translation: prior work finds that “prompt engineering is the key to improving translation performance” in low-resource scenarios (Khoboko et al., 2025). In line with this, we design structured prompts that guide our small Qwen models (0.5–3B parameters) to produce accurate medical translations. The prompt template includes an explicit instruction (e.g. “Translate the following sentence from Vietnamese to English: ...”) and may incorporate context cues (such as domain indicators or source/target language tags) to focus the model on the medical domain. Recent studies show that how we phrase the instruction – and whether we provide examples – can dramatically affect translation quality (Zhang et al., 2023). For instance, AFSP and other prompting frameworks demonstrate that providing well-chosen exemplar translations in a prompt (few-shot prompting) can help an LLM apply its knowledge more effectively (Tang et al., 2025). We leverage these insights by experimenting with multiple prompt styles as detailed below.

Prompt Templates We implement two prompt variants (see Appendix A) to compare different prompting strategies.

- **Instruction Prompt** (Appendix A.1): The prompt provides context by defining the translator’s role, required style, and target audience, then presents an input sentence for translation. This guidance helps the model apply

appropriate medical terminology and maintain a formal tone in its output.

- **Direct Prompt** (Appendix A.2): A direct command with no examples. This template simply tells the model the task and supplies the input sentence. Instruction-only prompts test the model’s zero-shot translation ability.

These templates allow us to compare performance with and without in-context examples. We choose the examples in the instruction prompt to be representative of medical language, since prior work has shown that the quality of example prompts strongly impacts results: using suboptimal or irrelevant examples can degrade translation performance (Zhang et al., 2023).

Each prompt template is applied uniformly during both fine-tuning and inference. That is, during fine-tuning we format every sentence pair into the chosen prompt structure, and the model learns to predict the translation as the prompt’s completion. During inference we use the same prompt format with new inputs. This consistency ensures the model has seen the instruction pattern during training, which is known to yield better performance (Khoboko et al., 2025).

3.2 Fine-tuning Strategy

Inspired of (Ding et al., 2021), we apply the bidirectional training strategy to fine-tune small language models. The motivation behind this approach comes from how humans learn foreign languages: by using translation examples in both directions (e.g., English to Vietnamese and Vietnamese to English) to better master bilingual knowledge. We employed both full-parameter fine-tuning and Low-Rank Adaptation (LoRA) to adapt the Qwen-3 1.7B model for bilingual machine translation tasks. The training data consisted of parallel corpora in English–Vietnamese (en–vi) and Vietnamese–English (vi–en).

Given a pre-trained model with parameters $\theta \in \mathbb{R}^{d \times d}$, the objective of fine-tuning is to minimize the negative log-likelihood (NLL) of the target sentence $Y = (y_1, y_2, \dots, y_T)$ conditioned on the source sentence $X = (x_1, x_2, \dots, x_S)$:

$$\mathcal{L}(\theta) = - \sum_{t=1}^T \log P_{\theta}(y_t \mid y_{<t}, X).$$

For full-parameter fine-tuning, all weights θ are updated according to gradient descent:

$$\theta' = \theta - \eta \nabla_{\theta} \mathcal{L}(\theta),$$

where η denotes the learning rate.

For LoRA, we preserve the original pre-trained weights and introduce trainable low-rank matrices $A, B \in \mathbb{R}^{d \times r}$ with $r \ll d$. The adapted weight matrix θ' is expressed as:

$$\theta' = \theta + \Delta\theta, \quad \Delta\theta = BA,$$

where $\Delta\theta$ is restricted to rank- r . Only A and B are optimized during training:

$$\begin{aligned} A' &= A - \eta \nabla_A \mathcal{L}(\theta + BA) \\ B' &= B - \eta \nabla_B \mathcal{L}(\theta + BA) \end{aligned}$$

This dual fine-tuning strategy leverages the capacity of full-parameter optimization while retaining the efficiency and memory benefits of LoRA. The combination ensures robust adaptation of the Qwen-3 1.7B model to English–Vietnamese bidirectional translation.

4 Experiments

4.1 Experimental Settings

For the Low-Rank Adaptation configuration, the experiments were conducted with the specific configuration set via the peft library. The rank (r) for the low-rank matrices was set to 64, with a corresponding alpha of 64. A dropout of 0 and a bias setting of "none" were used for optimal performance. LoRA was applied to the target linear modules. While fine-tuning all parameters, the training process was configured with a batch size of 4 and a gradient accumulation steps of 8, resulting in an effective batch size of 32. The optimizer used was adamw_8bit, with a learning rate of $2e-4$, a weight decay of 0.01, and a linear learning rate scheduler. A warmup of 5 steps was also included. All experimental setups were conducted on a NVIDIA A100 with 80G GPU.

4.2 Results and Discussion

Our submitted system, **Qwen3 1.7B (full fine-tune)**, achieves the best performance on the public test with BLEU scores of 48.60 for En–Vi and 36.10 for Vi–En (Table 1). Compared to the same base model tuned with LoRA, full fine-tuning yields a large gain for En–Vi (+10.21 BLEU, about 26.6%) and a minimal gain for Vi–En (+0.39 BLEU, about 1.1%). Against the strongest LoRA baseline in our set (Qwen2.5 3B), the full-finetuned

Table 1: BLEU scores of different models on the English–Vietnamese and Vietnamese–English translation tasks (evaluated on the public test).

Model Variants	Training Strategy	BLEU (en→vi)	BLEU (vi→en)
GPT-4o	Baseline	39.44	25.58
GPT-4.1	Baseline	41.47	27.89
Qwen2.5 0.5B	LoRA	37.99	26.25
Qwen2.5 1.5B	LoRA	42.39	27.25
Qwen2.5 3B	LoRA	44.37	27.95
Qwen3 0.6B	LoRA	32.51	22.41
Qwen3 1.7B	LoRA	38.39	35.71
Qwen3 1.7B	FPFT	48.60	36.10

Table 2: BLEU scores and corresponding rankings of different teams participating in both translation tasks (evaluated on the private test). Best scores are in bold, and second-best scores are underlined.

User	Ranking	BLEU score		Average
		Vi-En	En-Vi	
longday1102	1	25.4420	59.8540	42.6480
Tsunn	3	20.5962	43.2687	31.9325
cnv	4	11.2666	43.3784	27.3225
thindang (Ours)	2	<u>25.1085</u>	<u>50.6058</u>	<u>37.8571</u>

Qwen3 1.7B still leads by +4.23 BLEU on En–Vi. Relative to the GPT-4.1 baseline, our full-finetuned model improves by +7.13 BLEU (En–Vi) and +8.21 BLEU (Vi–En).

COMET¹ is an open-source evaluation framework developed by the Unbabel team, specifically designed for assessing the quality of Machine Translation (MT). It leverages pretrained models to generate semantic adequacy scores for source–translation pairs. In our experiments, we employed the Unbabel/wmt22-comet-da² model, which implements a reference-based regression approach built upon the XLM-R architecture. This choice enables robust semantic evaluation across both high- and low-resource language pairs. The evaluation was conducted on 200 samples drawn from the public dataset for both Vietnamese–English (Vi–En) and English–Vietnamese (En–Vi) translation tasks. The resulting system-level scores range from 0 to 1, with higher values indicating stronger semantic alignment. Table 3 presents the COMET scores obtained for each task.

The results indicate that most general content and straightforward phrases were accurately translated, Vietnamese grammar rules were generally respected, and many medical and technical terms

Table 3: Average COMET (wmt22-comet-da) evaluation scores for English–Vietnamese (En–Vi) and Vietnamese–English (Vi–En) translation tasks.

Task	Average Score (0–1)
En → Vi	0.822
Vi → En	0.811

were correctly preserved. However, several recurring issues were observed, including minor lexical or stylistic errors, occasional mistranslations of specialized terms, terminology inconsistencies, and critical omissions that reduce accuracy in some technical contexts. While these minor stylistic and phrasing issues slightly affect readability, they do not substantially hinder comprehension. A more detailed analysis of representative samples, illustrating these error types and their impact on translation quality, is provided in Appendix B.

To complement the automatic evaluations, we further report a small-scale human expert assessment conducted by the shared task organizers. This evaluation focused on adequacy and fluency in the medical domain, providing an additional perspective on translation quality beyond BLEU-based metrics. The results in Table 4 show that our system (*thindang*) achieved the highest overall human evaluation score, ranking first in both directions.

¹<https://github.com/Unbabel/COMET>

²<https://huggingface.co/Unbabel/wmt22-comet-da>

Table 4: Human evaluation scores in both translation tasks. Best scores are in bold

User	en-vi	vi-en	Average score	Ranking
thindang (Ours)	88.2	83.8	85.7	1
longday1102	87.7	80.4	84.0	2
cnv	84.1	78.0	81.0	3
Tsunm	74.4	79.9	77.1	4

These results indicate that the combination of domain-focused data, bidirectional formatting, and consistent instruction-style prompting during both fine-tuning and inference produces substantial gains, especially when the model undergoes complete parameter updates. The particularly large improvement in the En–Vi direction suggests that full fine-tuning more strongly benefits target side generation in Vietnamese (fluency and domain terminology), while the modest change in Vi–En implies either that LoRA already captured much of the reverse mapping ability or that the prompt effect is direction dependent.

Directional asymmetry (En–Vi vs Vi–En).

Across models, performance on En–Vi consistently exceeds Vi–En (mean gap about 12 BLEU). Possible causes include harder target-side generation into Vietnamese (morphology and idiom), tokenizer/pretraining bias toward English, dataset distributional imbalance, and BLEU’s sensitivity to morphological variation.

Final Ranking Table 2 presents the BLEU scores and rankings for four teams participating in a translation task. The scores are broken down into Vietnamese-to-English (Vi–En), English-to-Vietnamese (En–Vi), and an overall average. The team longday1102 achieved the top rank, securing the best scores in both translation directions, with a Vi–En score of 25.4420 and an En–Vi score of 59.8540, leading to a highest average score of 42.6480. The thindang (Ours) team performed commendably, earning the second-best ranking. This team achieved the second-highest scores in both tasks, with a Vi–En score of 25.1085 and an En–Vi score of 50.6058, resulting in a second-best average of 37.8571. Following these top two teams, Tsunn and cnv were ranked third and fourth, respectively, with significantly lower average BLEU scores of 31.9325 and 27.3225.

5 Conclusion

In this paper, we present our approach for the machine translation task of the VLSP 2025 shared task. The objective of this task was to apply pre-trained small language models, such as the Qwen 2.5 and Qwen 3 families, with limited parameters to a provided dataset for the medical domain. To address this challenge, we utilized Qwen 3 with 1.7B parameters as our foundation model, combined with a full-parameter fine-tuning approach. Compared to other versions and training strategies (e.g., LoRA), our solution achieved the best performance on the public test dataset. In the final ranking on the private test set³, our submitted model ranked second in terms of average BLEU score, as well as for the individual Vi–En and En–Vi translation tasks.

References

- Sheila Castilho and Rebecca Knowles. 2025. A survey of context in neural machine translation and its evaluation. *Natural Language Processing*, 31(4):986–1016.
- Liang Ding, Di Wu, and Dacheng Tao. 2021. [Improving neural machine translation by bidirectional training](#). In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 3278–3284, Online and Punta Cana, Dominican Republic. Association for Computational Linguistics.
- Long Doan, Linh The Nguyen, Nguyen Luong Tran, Thai Hoang, and Dat Quoc Nguyen. 2021. Phomt: A high-quality and large-scale benchmark dataset for vietnamese-english machine translation. *arXiv preprint arXiv:2110.12199*.
- Miguel Angel Rios Gaona, Raluca-Maria Chereji, Alina Secară, and Dragoş Ciobanu. 2023. Quality analysis of multilingual neural machine translation systems and reference test translations for the english-romanian language pair in the medical domain. In *Proceedings of the 24th Annual Conference of the European Association for Machine Translation*, pages 355–364.

³<https://aihub.ml/competitions/979#results>

- Lifeng Han, Serge Gladkoff, Gleb Erofeev, Irina Sorokina, Betty Galiano, and Goran Nenadic. 2024. Neural machine translation of clinical text: an empirical investigation into multilingual pre-trained language models and transfer-learning. *Frontiers in Digital Health*, 6:1211564.
- Pitso Walter Khoboko, Vukosi Marivate, and Joseph Sefara. 2025. Optimizing translation for low-resource languages: Efficient fine-tuning with custom prompt engineering in large language models. *Machine Learning with Applications*, 20:100649.
- Bryan Li, Mohammad Sadegh Rasooli, Ajay Patel, and Chris Callison-Burch. 2022. Multilingual bidirectional unsupervised translation through multilingual finetuning and back-translation. *arXiv preprint arXiv:2209.02821*.
- Raphaël Merx, Hanna Suominen, Lois Hong, Nick Thieberger, Trevor Cohn, and Ekaterina Vylomova. 2025. Tulun: Transparent and adaptable low-resource machine translation. *arXiv preprint arXiv:2505.18683*.
- C Ngo and 1 others. 2022. Mtet: multi-domain translation for english and vietnamese (2022). *arXiv preprint arXiv:2210.05610*.
- Long Nguyen, Tran Le, Huong Nguyen, Quynh Vo, Phong Nguyen, and Tho Quan. 2025. [Serving the underserved: Leveraging BARTBahnar language model for bahnaric-Vietnamese translation](#). In *Proceedings of the 1st Workshop on Language Models for Underserved Communities (LM4UC 2025)*, pages 32–41, Albuquerque, New Mexico. Association for Computational Linguistics.
- Nghia Luan Pham, Van Vinh Nguyen, and Thang Viet Pham. 2023. [A data augmentation method for english-vietnamese neural machine translation](#). *IEEE Access*, 11:28034–28044.
- Long Phan, Tai Dang, Hieu Tran, Trieu H Trinh, Vy Phan, Lam D Chau, and Minh-Thang Luong. 2022. Enriching biomedical knowledge for low-resource language through large-scale translation. *arXiv preprint arXiv:2210.05598*.
- Miguel Rios, Raluca-Maria Chereji, Alina Secara, and Dragos Ciobanu. 2022. Impact of domain-adapted multilingual neural machine translation in the medical domain. *arXiv preprint arXiv:2212.02143*.
- Lei Tang, Jinghui Qin, Wenxuan Ye, Hao Tan, and Zhi-jing Yang. 2025. Adaptive few-shot prompting for machine translation with pre-trained language models. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 39, pages 25255–25263.
- Nhu Vo, Dat Quoc Nguyen, Dung D Le, Massimo Piccardi, and Wray Buntine. 2024. Improving vietnamese-english medical machine translation. *arXiv preprint arXiv:2403.19161*.
- Biao Zhang, Barry Haddow, and Alexandra Birch. 2023. Prompting large language model for machine translation: A case study. In *International Conference on Machine Learning*, pages 41092–41110. PMLR.

A Appendix: Example of Instruction and Direct Prompts

A.1 Instruction Prompt Example

Instruction Prompt

Below is an instruction that describes a task, paired with an input that provides further context. Write a response that appropriately completes the request.

Instruction:

You are a highly skilled medical translator, an expert in translating complex medical texts from English to Vietnamese for a professional audience of doctors and researchers. Your translations must be precise, use standard Vietnamese medical terminology, and maintain a formal, objective tone. Output only the finished translation.

Translate the following English medical text to Vietnamese.

Input:

The patient was diagnosed with acute bronchitis and prescribed antibiotics.

Response:

Bệnh nhân được chẩn đoán bị viêm phế quản cấp tính và được kê đơn thuốc kháng sinh.

A.2 Direct Prompt Example

Direct Prompt EN → VI

<|im_start|>user

You are a highly skilled medical translator.

Translate the following English medical text to Vietnamese:

The patient was diagnosed with acute bronchitis and prescribed antibiotics.

<|im_end|>

<|im_start|>assistant

<think>

</think>

Bệnh nhân được chẩn đoán bị viêm phế quản cấp tính và được kê đơn thuốc kháng sinh<|im_end|>

Direct Prompt VI → EN

<|im_start|>user

You are a highly skilled medical translator.

Translate the following Vietnamese medical text to English:

Bệnh nhân được chẩn đoán bị viêm phế quản cấp tính và được kê đơn thuốc kháng sinh.

<|im_end|>

<|im_start|>assistant

<think>

</think>

The patient was diagnosed with acute bronchitis and prescribed antibiotics.<|im_end|>

B Detailed Analysis

Sample 1

English Source: Results: Percentage of card's holders who knew the finance-free utilization of the first registered public health services was 44.5% and being provided health insurance information was 34.8%.

Vietnamese Translation: Kết quả: Tỷ lệ người dân biết đến dịch vụ y tế công không cần thẻ BHYT là 44,5% và được cung cấp thông tin về BHYT là 34,8%.

Reference Translation: Kết quả: Tỷ lệ người biết được khám chữa bệnh (KCB) miễn phí tại nơi đăng ký ban đầu chiếm 44,5%, được cung cấp thông tin về bảo hiểm y tế (BHYT) chiếm 34,8%.

Analysis: The phrase "dịch vụ y tế công không cần thẻ BHYT" misrepresents the concept of "finance-free utilization of the first registered public health services."

Sample 2

English Source: To report the clinical and radiological outcomes after corrective osteotomy in ankylosing spondylitis patients.

Vietnamese Translation: Nhận xét kết quả điều trị phẫu thuật gãy xương chính mũi do viêm cột sống dính khớp.

Reference Translation: Đánh giá kết quả về lâm sàng và hình ảnh học X-quang sau điều trị phẫu thuật cắt chêm xương đốt sống để nắn chỉnh gù toàn bộ cột sống do bệnh viêm cột sống dính khớp (VCSDK).

Analysis: Critical mistranslation of "corrective osteotomy" as "gãy xương chính mũi" (nasal fracture).

Sample 3

English Source: Patients who came to the examination rooms of the Out-Patient Examination Department of Nhat Tan hospital with hypertension that the doctor thought about the white coat effect during 8 months from October 1, 2020 to May 31, 2021.

Vietnamese Translation: Bệnh nhân đến khám tại phòng khám Khoa Khám bệnh của bệnh viện Nhat Tân có tăng huyết áp mà bác sĩ nghĩ đến hiệu ứng áo choàng trắng trong thời gian 8 tháng từ 01/10/2020 đến 31/05/2021.

Analysis: Awkward phrasing reduces readability; minor semantic shifts, but overall meaning is preserved.

Sample 4

English Source: From 09/2020 - 09/2021, the patient had upper GI endoscopy (06 times) and biopsies (06 times).

Vietnamese Translation: Từ tháng 09/2020 – 09/2021, bệnh nhân được nội soi tiêu hoá trên (06 lần) và sinh thiết (06 lần).

Analysis: Translation is understandable but less precise than reference ("nội soi dạ dày" would be more specific).

Sample 5

Vietnamese Source: Phương pháp: Thiết kế nghiên cứu mô tả cắt ngang được thực hiện trên 928 người trưởng thành có thẻ bảo hiểm y tế tại 2 huyện Phone Hong và Keo Oudom, tỉnh Viêng Chăn.

Translation: Methods: A cross-sectional study was conducted on 928 adult health insurance card's holders in two districts of Phone Hong and Keo Oudom, Vinh Chau town, Yen Bai province.

Analysis: Critical mistranslation of geographical regions leads to significant confusion. Despite fluent language, the geographic error lowers accuracy.