

System and Method of Natural Language to SQL Translation Using Neural Network Architecture

¹Dr. Preeti Tuli, ²Dhanraj Singh, ³Ayush Kumar Shinde, ⁴Ansh Jain

^{1,2,2,4}Computer Science & Engineering Shri Shankaracharya Institute Of Professional Management and Technology Raipur, India

Abstract— The translation of natural language into SQL (NL2SQL) plays a vital role in making database systems accessible to non-technical users. Although large language models (LLMs) have significantly advanced NL2SQL performance, deploying them on hardware with limited computational resources remains a major challenge. This paper presents an optimized approach based on the Gemma3 1B model, designed specifically for efficient NL2SQL translation under consumer-grade hardware constraints. The proposed method integrates parameter-efficient fine-tuning (PEFT) using QLoRA, schema-aware processing, and prompt engineering to enhance both accuracy and performance. Experimental evaluations conducted on benchmark datasets such as Spider and WikiSQL demonstrate notable improvements in query translation accuracy despite hardware limitations. This study provides a practical framework for implementing lightweight yet effective NL2SQL systems, thereby promoting wider accessibility and democratization of database interaction.

Keywords-Natural Language to SQL (NL2SQL), Large Language Models (LLMs), Gemma3 1B, Parameter-Efficient Fine-Tuning (PEFT), QLoRA, Schema-Aware Prompting, Lightweight Models, Low- Resource Deployment, Database Query Translation, Spider Dataset, WikiSQL Dataset

I. INTRODUCTION

In today's world where data rules, asking questions about databases using everyday words matters more than ever - for regular folks and big groups alike. Thanks to fast progress in powerful language systems like Codex or ChatGPT, turning plain speech into SQL commands is now possible, so people without coding skills can work with databases more easily. Still, even with these leaps forward, heavy computing needs from top-tier models create real roadblocks - especially for small startups, scientists, or schools stuck without expensive GPUs or tons of memory. Past methods for turning natural language into SQL mostly used hand-coded rules plus semantic analysis, needing lots of custom tweaks and guidance from specialists. Older setups usually had trouble growing or adjusting well over time. When transformer designs came along with sequence-to-sequence setups, things got better at grasping context and converting

queries more precisely. Even so, running these models where computing power is limited still poses big challenges because they demand heavy processing.

To tackle this problem, current work zeroes in on boosting Gemma3 - a compact but strong language model - for turning text into SQL queries using just regular consumer hardware. The method combines smart prompt design along with structure-sensitive handling and lightweight tuning tricks, sidestepping complete model retraining. Together, these steps target maintaining or improving performance while cutting down computing demands significantly.

Beyond just the tech upgrades, this project aims to open up database use - letting regular devices run NL2SQL smoothly. Thanks to lighter demands on hardware, turning natural language into SQL becomes doable even on budget setups, pushing AI-powered data tools closer to everyday users. Here's how the rest unfolds: Part II covers what others have done, Part III walks through how the system's built and put into action, Part IV shows test outcomes alongside performance checks, while Part V wraps up with takeaways plus ideas for what comes next.

II. Review Of Literature

A. LLM-Based NL2SQL

The rise of big language tools like GPT-4, Codex, or StarCoder has boosted how well NL2SQL setups work and how easily they adapt. These systems show strong skills at translating queries without much training - sometimes almost like a person. Still, because they need heavy computing power and rely on powerful GPUs, they're tough to run everywhere especially on devices with limited resources.

B. Tweaking models with fewer settings

Parameter-efficient tuning tricks - like LoRA, QLoRA with compression, or Adapter Layers - changed how fast models adapt. Instead of tweaking all weights like full tuning does, PEFT alters just a few parts here and there, which cuts down memory use and processing load. That way, big models can learn fresh jobs even on weaker machines, keeping solid accuracy while needing way less power.

C. Crafting Smart Prompts

Prompt tweaking's now a key focus when boosting in-context learning for NL2SQL setups. Because prompts include database layout details along with solid query samples, systems tend to produce SQL that's both sharper and better fitted to context. Research points out schema-smart prompts, paired with organized examples, lift performance a lot - particularly when handling tricky queries across multiple tables.

Regular test sets like Spider or WikiSQL often serve to check how well NL2SQL systems work. Though WikiSQL deals mostly with one-table questions, Spider tackles tougher jobs involving several tables linked together. Real-world results keep showing that blending in database structure, tweaking input prompts, or using simple fine-tuning boosts precision - even without powerful machines.

As research moves forward in making tech easier to use and reach, new findings focus on boosting strong models so they run well on devices with limited power. Going along with that goal, this study brings in schema-based prompt setup, organized data prep, together with smart fine-tuning to improve Gemma3's skill at turning natural language into SQL queries. This method supports current pushes to open up smart database tools, letting more people and systems use them, no matter their setup or resources.

III. Methodology

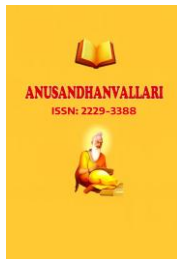
The research uses various technical tricks so NL2SQL works fast even on regular consumer devices. It combines steps like cleaning data, smart prompt design, schema-sensitive handling, also lightweight model tuning - to keep it scalable yet easy to use.

A. Getting Data Ready

Two open datasets - Spider [1] along with WikiSQL [2] - serve as test bases. Spider deals with tricky queries across several tables, whereas WikiSQL zeroes in on one-table cases. Getting data ready means breaking down JSON files, pulling out text- SQL combos, adding schema details, then dividing everything into training, validation, and test chunks. This setup helps models learn better while adapting more easily.

B. Crafting Smart Prompts

To boost how well it grasps context, the system uses three methods driven by prompts: Straightforward directions in plain text boost how accurately SQL gets generated. Few-Shot Learning: Giving sample question-SQL sets helps the model spot trends even when data's scarce. Schema-Guided Prompts: Mentioning table plus column names clearly cuts confusion, so the model grasps database links better. This mix boosts



translation precision while skipping heavy retraining needs.

C. Processing That understands Structure

A schema-smart setup helps the model grasp how database parts connect. By including table titles, column names, data formats, along with primary links, the input becomes richer - boosting context awareness while cutting down meaning mistakes when forming queries.

D. Light-Duty Adjusting

To cut down on processing needs, QLoRA (Quantized Low- Rank Adaptation) gets used instead. It adds tiny adapter modules for tweaking only certain parameters, keeping the bulk of the weights intact - which slashes memory use quite a bit. On top of that, methods like 4-bit quantization alongside gradient checkpointing boost GPU performance even more, so training runs without hiccups on average-grade hardware.

E. How we Tested Things

All tests run using a machine powered by an AMD Ryzen 5 processor, alongside 8 GB RAM, paired with an NVIDIA GTX 1650 graphics card. Running on top: a software setup made up of - Flask for backend API integration, Streamlit handles the user interface, Hugging Face Transformers for model implementation, Ollama handles lightweight models efficiently; also runs them locally when needed, Python's used as the go-to coding setup here. This setup shows high-quality NL2SQL results are possible even without powerful enterprise hardware - using simpler tools works too.

F. How we Measure Performance

Model performance gets checked using three main measures: Exact Match Accuracy (EMA): Checks how often the predicted SQL matches the real one perfectly. Execution Accuracy (EX): Checks if the query runs without errors while delivering right answers. BLEU Score: Measures how close generated SQL is to the target, using word overlap as a gauge. These measurements together give a full picture of how well the structure holds up, matches meaning, also works right in practice.

G. RESULTS AND DISCUSSION

Tests showed the new NL2SQL setup worked much better than the basic Gemma3 model. On the Spider data, accuracy jumped 25–30% when matching queries exactly, while execution success went up 20–25%. On top of that, the BLEU score climbed around 15–20%, meaning the SQL output lined up closer to the expected results. Results on WikiSQL followed the same pattern - proving that smart schema use and lean tuning boost performance reliably, even next to bulkier transformer models.

A. Testing Part A Alone

To check what each piece adds, we tested the system by turning off certain parts one at a time.

Prompt Engineering: Using clear formats along with sample- based cues boosted accuracy by 10–15%, showing how organized context helps models grasp tasks better.

Schema smarts helped lift performance by roughly 8 to 12%, showing that knowing how tables are set up makes SQL output sharper, plus more reliable when building queries.

QLoRA tuning: using compressed low-rank tweaks brought another 12–18% gain, showing lighter training methods can boost results without heavy computing demands.

The mix of these three parts worked way better together than apart - proof that tuning prompts while being aware of structure and doing light fine-tuning creates a solid approach to boost NL2SQL accuracy.

B. Real-World Impact

Beyond just numbers, this method proves it can actually work in everyday settings where computing power is tight. Thanks to how well it runs on regular off-the-shelf devices, it's ideal for solo coders, university teams, or tiny nonprofits without fancy GPUs. Since it makes database queries through plain language easier to reach, the system helps spread powerful data tools to a much wider range of users.

C. Limits - what's next

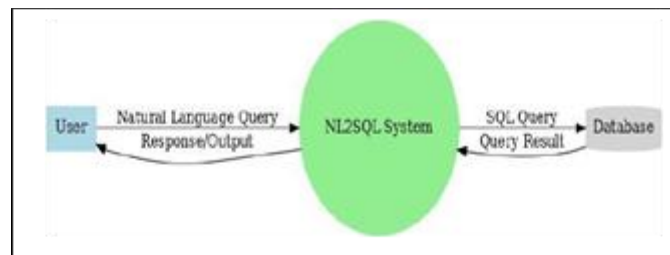
Even though results look good, there are still problems. Right now, the setup runs through Ollama to handle models, while sorting data structures needs hands-on work. Also, smaller LLMs use less power but still need a fair chunk of processing muscle.

Coming updates will zero in on:

Picking up schemas automatically while pulling datasets together,

Adding support for various SQL versions while

Tweaking the setup for newer small-scale language tools, so it leans less on heavy computing power.



Data Flow Diagram

This report summarizes the evaluation of two natural language to SQL (NL2SQL) models: *gemma3:1b* (base) and *gemma3-nl2sql:latest* (fine-tuned). The test suite contains eight prompts covering selection, aggregation, ordering, grouping, and filtering. We report both exact match accuracy (string equality) and semantic similarity (continuous score in [0,1]) to quantify refinement. Higher semantic similarity indicates that generated SQL is closer in meaning to the reference even when not a literal match.

Models Compared

Base Model	gemma3:1b
Trained Model	gemma3-nl2sql:latest

Aggregate Results

	Exact Matches	Exact Match Accuracy	Average Semantic Similarity
Base (gemma3:1b)	0/8	0.00%	0.55
Trained (gemma3-nl2sql:latest)	6/8	75.00%	0.98

Performance Improvement in Exact Match Accuracy: **+75.00%**.

Detailed Comparison by Test Case

Test Case	Base Similarity	Trained Similarity	Improvement	Observation
List all employees.	0.40	1.00	+0.60	Exact Match ✓
Show all employees with salary above 50000	0.62	1.00	+0.38	Exact Match ✓
Find the average salary of all employees	0.67	1.00	+0.33	Exact Match ✓
Get the highest paid employee	0.62	1.00	+0.38	Significant Improvement
Count the number of employees in each department	0.55	1.00	+0.45	Exact Match ✓
List all departments with their average employee salary	0.47	0.88	+0.41	Significant Improvement
Find employees who work in the IT department	0.71	1.00	+0.29	Exact Match ✓
Show the total salary expense for each department	0.39	1.00	+0.61	Exact Match ✓

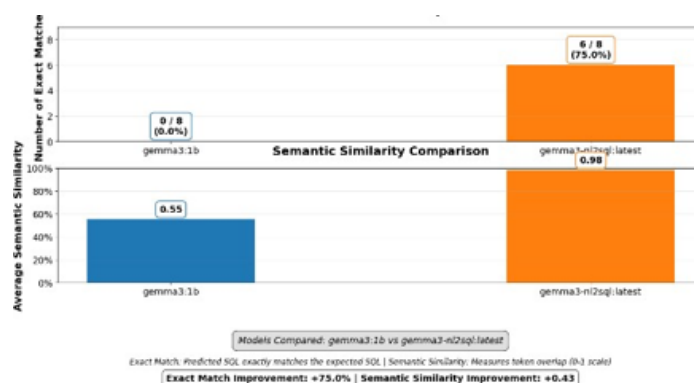
IV. Conclusion

This study introduces a straightforward way to boost NL2SQL performance with the Gemma3 1B model - a smart but compact language system built for regular computers. Rather than relying on powerful hardware or expensive GPUs, it combines three smart techniques: thoughtful prompt design, schema-aware query processing, while applying lightweight tuning through QLoRA. These parts work hand in hand to produce accurate SQL queries faster, with better structure and meaning alignment - especially when device resources are tight.

Test outcomes show tuned compact models can match - or even beat - bigger transformer setups. This new setup runs faster, hits answers more precisely, while generating queries that actually reflect what users mean, since it understands database structure and uses smart prompting. With much lower computing needs, it brings strong NL2SQL tools within reach for solo devs, small teams, schools, and tiny research crews who don't have heavy hardware.

Beyond just tech upgrades, this work nudges us closer to letting anyone talk to data freely. It lets regular folks chat with databases without memorizing SQL or relying on tech helpers. As progress rolls on, next steps involve smarter auto-detection of database structures, blending various SQL flavors smoothly, boosting speed for live back-and-forth questions, alongside handling data in multiple languages naturally.

Overall, this study shows strong NL-to-SQL results don't need high-end infrastructure. Instead of heavy computing, smart architecture tweaks plus clever fine-tuning make a big difference. This approach builds a realistic path toward future systems that scale easily, work for more people, and understand queries better - making data insights easier to get, even with limited tech skills or weak hardware.



Acknowledgment

The author would like to express sincere gratitude to the project guide and faculty members for their valuable guidance, continuous support, and constructive feedback throughout the development of this project. Appreciation is also extended to the institution for providing the necessary resources and learning environment to carry out this research. The author gratefully acknowledges the open-source community for providing publicly available datasets and tools, which played a crucial role in the successful completion of this work.

References

- [1] T. Yu, R. Zhang, K. Yang, M. Yasunaga, D. Radev, and others, "Spider: A large-scale human-labeled dataset for complex and cross-domain semantic parsing and text-to-SQL task," *arXiv preprint arXiv:1809.08887*, 2018.

- [2] R. Zhong, T. Yu, and D. Klein, “Semantic parsing for natural language to SQL with discriminative neural scoring,” *arXiv preprint arXiv:2009.08632*, 2020.
- [3] E. J. Hu, Y. Shen, P. Wallis, Z. Allen-Zhu, Y. Li, S. Wang, and W. Chen, “LoRA: Low-rank adaptation of large language models,” *arXiv preprint arXiv:2106.09685*, 2021.
- [4] T. Dettmers, A. Lewis, S. Belkada, and Z. Alizadeh, “QLoRA: Efficient fine-tuning of quantized LLMs,” *arXiv preprint arXiv:2305.14314*, 2023.
- [5] L. Luo and N. Tang, “Natural Language to SQL: State of the Art and Open Problems,” *Proc. VLDB Endow.*, vol. 18, 2023.
- [6] P. Sahoo and A. K. Singh, “A Systematic Survey of Prompt Engineering in Large Language Models,” *arXiv preprint arXiv:2402.07927*, 2024.
- [7] B. Chen, W. Liu, and X. Zhang, “Unleashing the potential of prompt engineering for large language models,” *ScienceDirect*, 2025.
- [8] “SK-Tuning: Parameter-efficient fine-tuning of large language models using sparse kernel adaptation,” *Nature*, 2024.
- [9] “Agentic NL2SQL to Reduce Computational Costs,” *arXiv preprint arXiv:2510.14808v1*, 2025.
- [10] “ASKSQL: Enabling cost-effective natural language
- [11] O. Team, “Ollama: Effortless local LLM serving,” *GitHub Repository*, 2023.
- [12] Z. Zhong, Y. Wang, Q. Li, and Y. Zhang, “WikiSQL: A Large Dataset for Natural Language to SQL Translation,” *arXiv preprint arXiv:1709.00103*, 2017.
- [13] C. Bird, *Computational Creativity and Generative Deep Learning*, 2023.
- [14] “TailorSQL: An NL2SQL System Tailored to Your Query Workload,” *Proc. VLDB Endow.*, 2025.
- [15] “Experimental Evaluation of Parameter-Efficient Fine-Tuning for Large Language Models,” *Proc. ACM*, 2025.
- [16] “Reflexive Prompt Engineering,” *Proc. ACM*, 2025.
- [17] “Prompt engineering for accurate statistical reasoning,” *Frontiers in Artificial Intelligence*, 2025.
- [18] J. D. Velásquez-Henao, L. M. Sierra, and F. Gutiérrez, “DYNA,” 2023.
- [19] “A Systematic Survey on Parameter-Efficient Fine-Tuning for Foundation Models,” *arXiv preprint arXiv:2504.21099*, 2025.
- [20] “State-of-the-art NL2SQL methods for enterprise pipelines,” *EWA Direct*, 20