

2026 A Real-Time LLM-Driven Rushdown and Zoning Specialist for DareFightingICE

ANJAL SHRESTHA

SI THU AUNG

Team Iconic

August 7, 2026

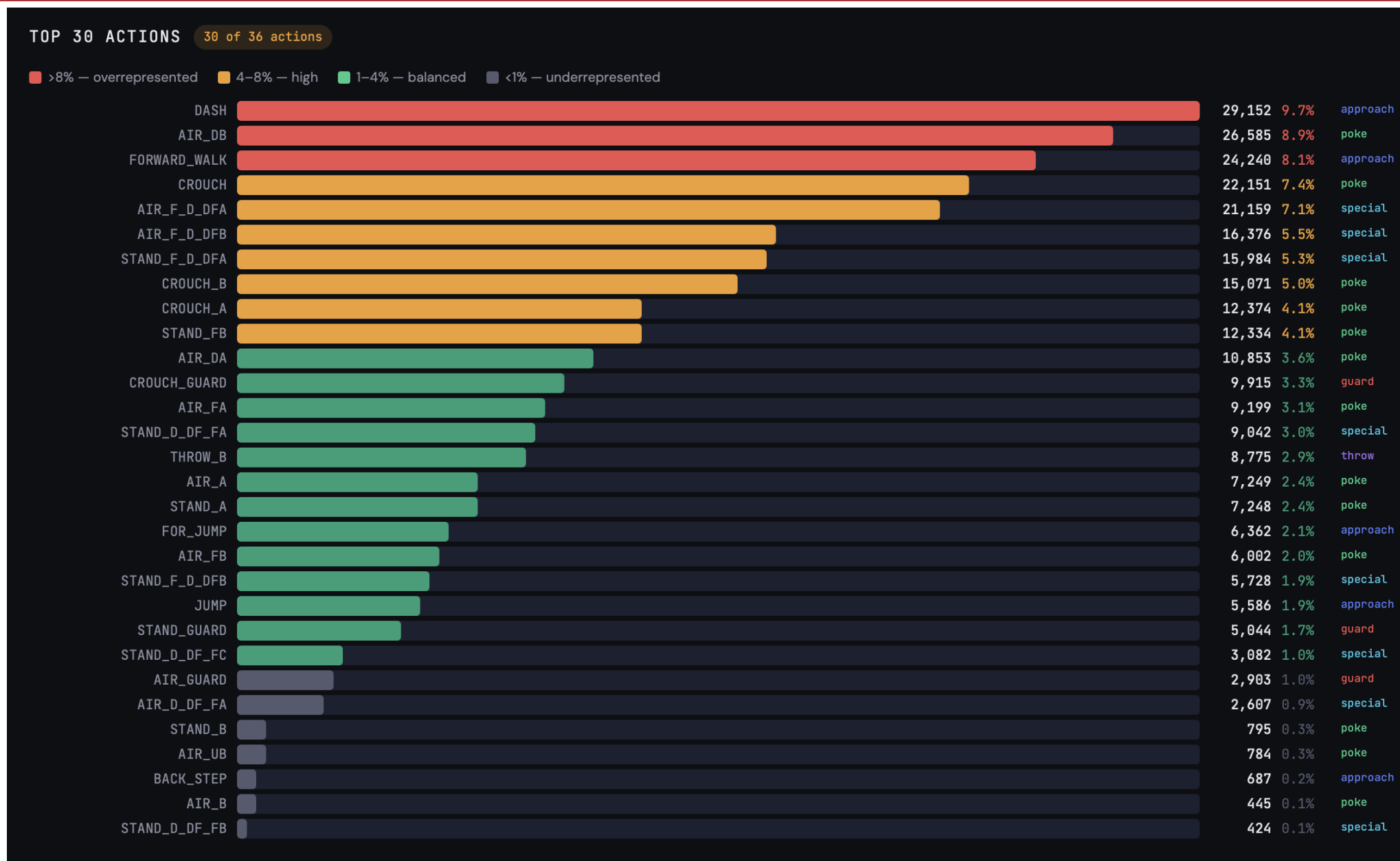
Intelligent Computer Entertainment Lab
College of Information Science and Engineering
Ritsumeikan University

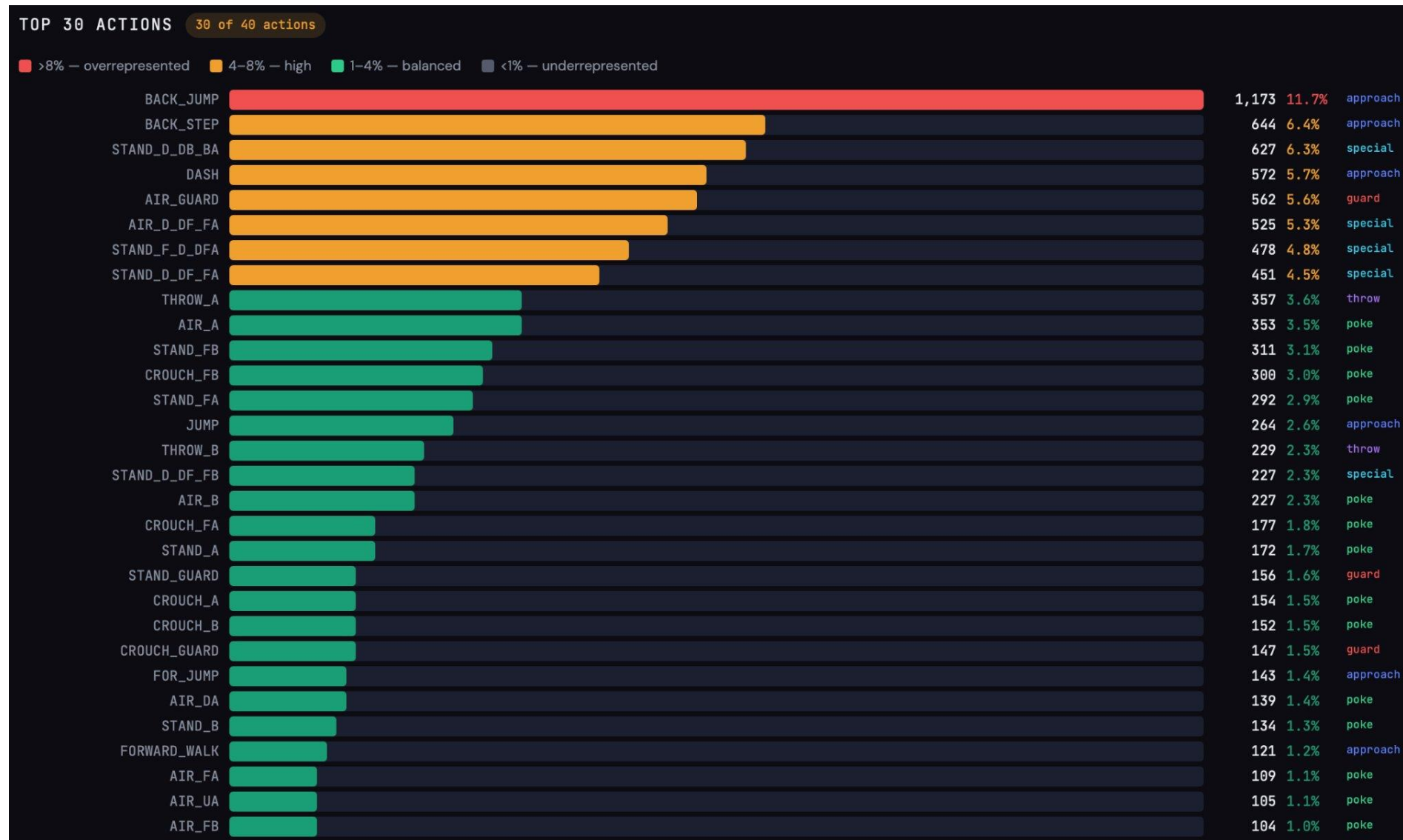
Main Changes for midterm

- Changed the model to Qwen1.5 with 1.5 billion parameters because it is lightweight enough to process a frame instantly and smart enough to understand complex variables.
- Changed the data structure for the rushdown agent so it retains a list of its previous moves.
- Combined three data set from hugging face for zoning.
- Modified agent.py and prompt.py so the trained model receives data in the exact structure it was trained on.
- Maintained the original data structure for the zoning agent, using the same training script except for the data feeding method, which matches the provided sample.

- **1-GPU Optimization:** Altered the training script to run perfectly on a single GPU by using a `per_device_batch_size` of 16 and a `gradient_accumulation` of 4 to safely saturate VRAM without crashing.
- **Training Length:** 3 full Epochs (roughly 11,500 steps) for zoning and 5 full Epochs(8,000) for zoning.
- **Checkpoint Strategy:** Standard safety backups trigger every 500 steps (keeping only the last 2 to save hard drive space).
- Custom callbacks trigger at the end of Epoch (1 to 5) to export permanent, playable models.
- **Label Masking (-100):** Implemented a crucial code change so the AI only calculates loss on the *predicted action* and ignores the instruction prompt. This prevents the bot from hallucinating or repeating text mid-fight.

Data For Rushdown





- **The Logic Behind the Data:** Raw gameplay data structure was too complex, model training on previous structure had to think like distance position and had too many things when training, so by reduction the data complexity it gives us more time to think for model which is crucial for the Rushdown model.
- **Why Not Another Model?** First used *Qwen-3B*: Better tactical reasoning, but completely fails the 20ms real-time latency test and longer training time.
- **Qwen-0.5B:** Lightning fast, but struggles with nuanced state-awareness (e.g., it might forget to block when low on health). 1.5B provides the necessary balance.

- Real-Time LLM Viability:** Proves that lightweight models like Qwen-1.5B can successfully run inside a strict 20ms real-time loop without causing game lag.
- Data Balancing Shapes Identity:** Demonstrates that starving out passive actions is just as critical as feeding the model optimal combos to force a hyper-aggressive tactical playstyle.
- Scalable Architecture Blueprint:** Establishes a robust, corruption-proof training pipeline that can easily be reused to train distinct models for any future combat archetype.

Thank You



Ritsumeikan Intelligent Computer Entertainment