

Multi-Task Deep Learning for Joint Mobility Management and Resource Allocation in 5G Heterogeneous Dense Networks

Motivation & Contributions

- Handover (HO) and resource allocation (RA) are strongly coupled in dense 5G networks.
- Single-task optimisation ignores shared RAN state information and can lead to suboptimal decisions.

Key contributions

- Proposed a unified multi-task learning framework for joint HO and RA decisions.
- Using shared autoencoder with task specific heads for robust latent representations and network prediction.
- Parameter-efficient design with better system-level throughput and delay.

System Model

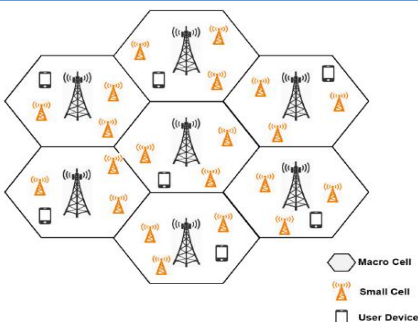
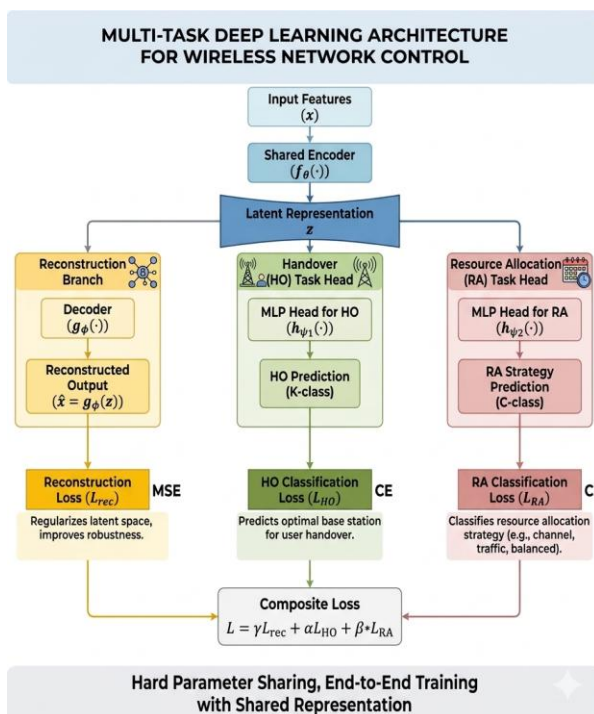


FIGURE 1: System architecture

- Heterogeneous dense network
- Mobile users
- mmWave small cells (24 GHz)
- Macro cells (3.8 GHz)

Proposed MTL Framework

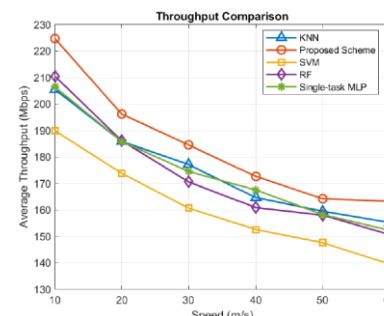


- Shared encoder learns a compact latent representation from network such as RSRP, SINR, load/utilisation and traffic demand.
- Two task-specific MLP heads predict HO mobility and RA strategy in parallel.
- A reconstruction decoder regularises the latent space.

Experimental Setup

- Dataset generated from MATLAB-based system simulation
- 100 mobile users with diverse trajectories
- Data split: 70% Training, 10% Valid, 20% Test
- Strong baselines: SVM, RF, KNN, STL-MLP

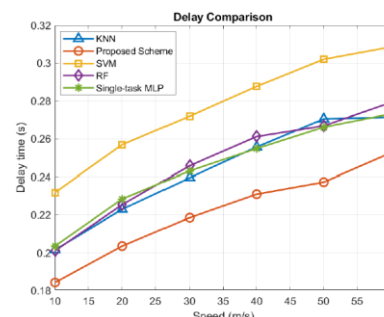
Key Results



15.6% ↑

avg. network throughput

Average throughput comparison

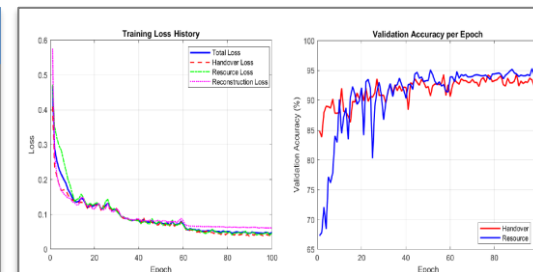


20.2% ↓

delay reduction

Delay time comparison

- Joint learning improves throughput and reduces delay under mobility.



Our Model	HO task	94.15%	0.862	0.928	0.894
Our Model	RA task	93.25%	0.946	0.904	0.925

- Training loss:** rapid convergence in early epoch + reconstruction loss improves robustness
- Validation accuracy:** fast accuracy improvement + Good generalization on validation set
- Prediction accuracy:** HO accuracy: 94.15%, RA accuracy: 93.25%, good F1 score

Conclusion

- MTL jointly optimises HO and RA from a unified RAN state representation.
- The shared AE-regularized latent space improves robustness and coordination across tasks.
- The proposed method achieves higher throughput and lower delay with efficient inference.