

# Weight Uncertainty in Neural Networks

Yufan Wang (yw575), Max Bronckers (mojb2), Alan Clark (ajc348)

## Motivation

- Regular feed forward neural networks use **point estimates** of weights

|               |                      |                                |                |
|---------------|----------------------|--------------------------------|----------------|
| Common Issues | Prone to Overfitting | Classic Regularisation Methods | Early Stopping |
|               | Overly Confident     |                                | Weight Decay   |
|               | Non-Robust           |                                | Dropout        |

## Bayes by Backprop (BBB)

- Instead of point estimates, **learn probability distribution on the weights**

|                |                 |             |
|----------------|-----------------|-------------|
| Regularisation | Model Averaging | Exploration |
|----------------|-----------------|-------------|

- Exact Bayesian inference is intractable, instead use variational approximation to the posterior
- Find parameters  $\theta$  of a distribution on the weights  $q(\mathbf{w}|\theta)$

$$\begin{aligned}
 \theta^* &= \arg \min_{\theta} \mathcal{F}(\mathcal{D}, \theta) = \arg \min_{\theta} \underbrace{\text{KL}[q(\mathbf{w}|\theta) || P(\mathbf{w})]}_{\text{Complexity Cost}} - \underbrace{\mathbb{E}_{q(\mathbf{w}|\theta)}[\log P(\mathcal{D}|\mathbf{w})]}_{\text{Likelihood Cost}} \\
 &= \arg \min_{\theta} \int q(\mathbf{w}|\theta) [\log q(\mathbf{w}|\theta) - \log P(\mathbf{w})P(\mathcal{D}|\mathbf{w})] d\mathbf{w} \\
 &= \arg \min_{\theta} \mathbb{E}_{q(\mathbf{w}|\theta)} [f(\mathbf{w}, \theta)]
 \end{aligned}$$

- Generalisation of the Gaussian reparameterisation trick

$$\frac{\partial}{\partial \theta} \mathbb{E}_{q(\mathbf{w}|\theta)} [f(\mathbf{w}, \theta)] = \mathbb{E} \left[ \frac{\partial f(\mathbf{w}, \theta)}{\partial \mathbf{w}} \frac{\partial \mathbf{w}}{\partial \theta} + \frac{\partial f(\mathbf{w}, \theta)}{\partial \theta} \right]$$

- Approximation of the exact cost using **Monte Carlo sampling**

$$\mathcal{F}(\mathcal{D}, \theta) \approx \sum_{i=1}^n \log q(\mathbf{w}^{(i)}|\theta) - \log P(\mathbf{w}^{(i)}) - \log P(\mathcal{D}|\mathbf{w}^{(i)})$$

- Inference performed using an **uncountably infinite neural networks**

$$P(\hat{\mathbf{y}}|\hat{\mathbf{x}}) = \mathbb{E}_{P(\mathbf{w}|\mathcal{D})} [P(\hat{\mathbf{y}}|\hat{\mathbf{x}}, \mathbf{w})]$$

## Reparameterisation Tricks

$$w_{i,j} = \mu_{i,j} + \sigma_{i,j} \epsilon_{i,j}$$

$$\epsilon_{i,j} \sim \mathcal{N}(0, 1)$$

Blundell et al. [1]

$$b_{m,j} = \gamma_{m,j} + \sqrt{\delta_{i,j}} \zeta_{m,j}$$

$$\gamma_{m,j} = \sum_{i=1} a_{m,i} \mu_{i,i} \quad \delta_{m,j} = \sum_{i=1} a_{m,i}^2 \sigma_{i,i}^2 \quad \zeta_{m,j} \sim \mathcal{N}(0, 1)$$

Kingma et al. [2]

## Regression

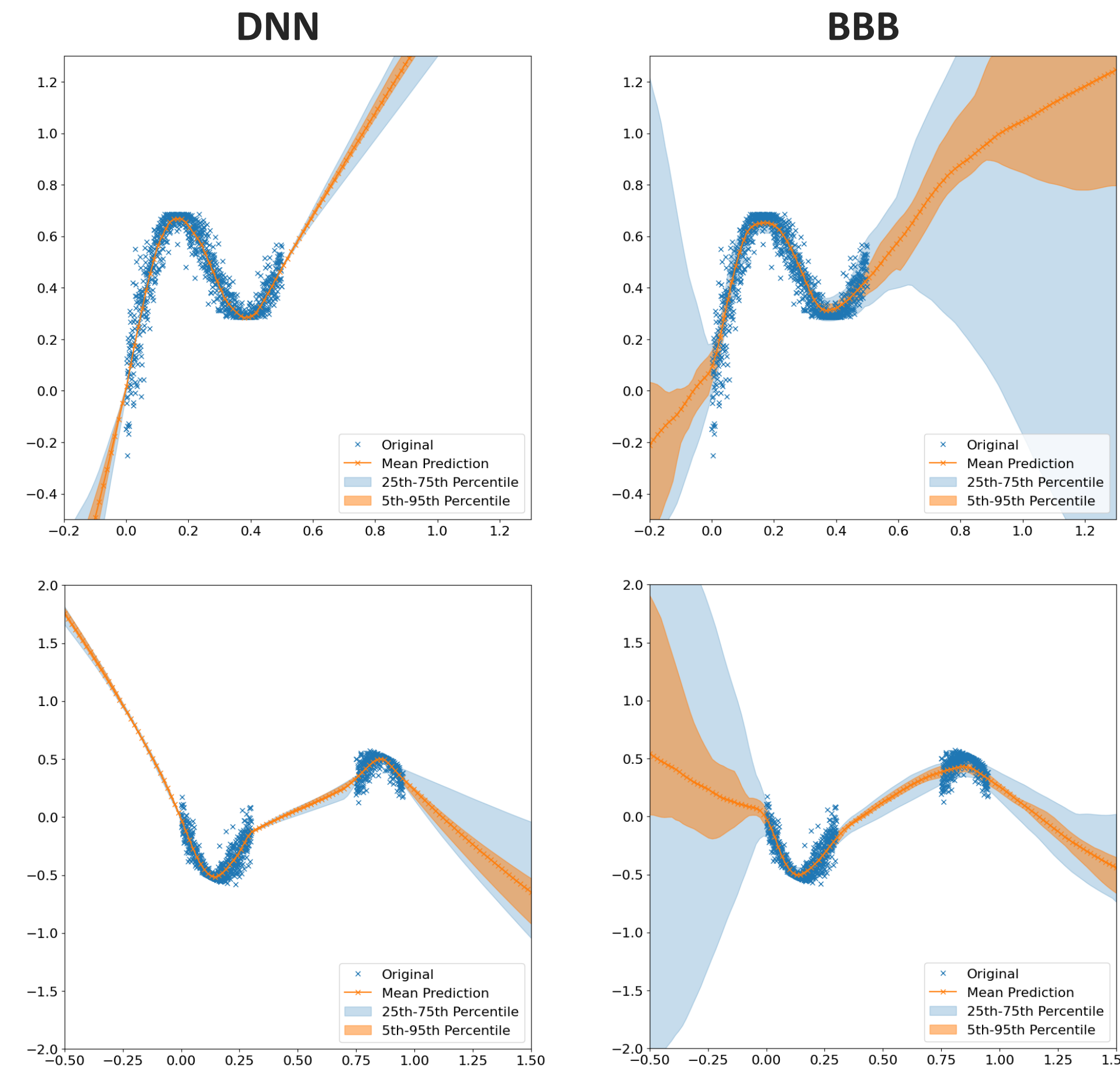


Figure 1. Regression against synthetic data using a regular DNN and BBB

- BBB exhibits significant uncertainty outside observed domain
- DNN exhibits characteristic overconfidence

## MNIST

Table 1. Classification errors

| Hidden Units | DNN    | DNN (Dropout) | BNN           |
|--------------|--------|---------------|---------------|
| 400          | 2.04 % | <b>1.33 %</b> | 1.44 %        |
| 800          | 1.68 % | <b>1.41 %</b> | 1.50 %        |
| 1200         | 1.66 % | 1.52 %        | <b>1.36 %</b> |

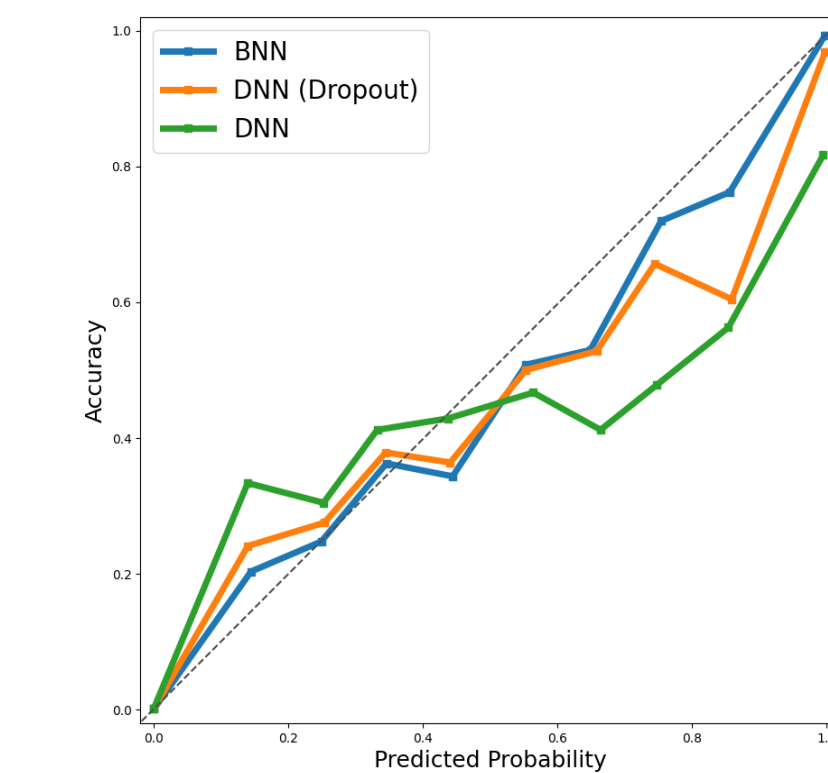


Figure 2. Reliability diagram

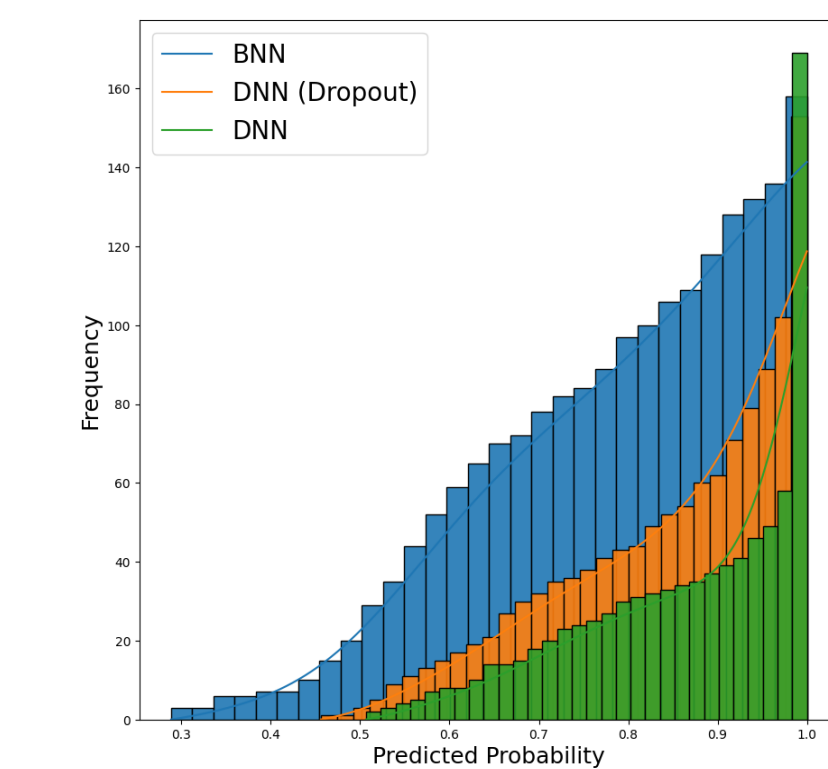


Figure 3. Cumulative distribution of confidence scores for incorrect predictions

- BBB is **more calibrated**
- Lower degree of confidence in its predictions

- DNN exhibits overconfidence in its incorrect predictions

## Contextual Bandits

- Dataset: UCI Mushroom Bandit
- Thompson sampling: picking an action trades-off between exploitation and exploration
- BNN expected to exhibit systematic exploration

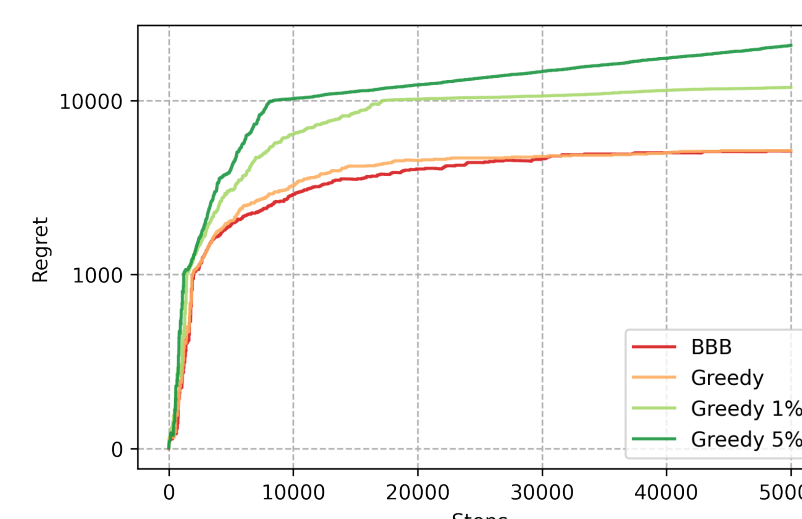


Figure 4. Comparison of cumulative regret for various agents on the mushroom bandit task (step size = 5000, gamma = 0.5, lr = 1e-4)

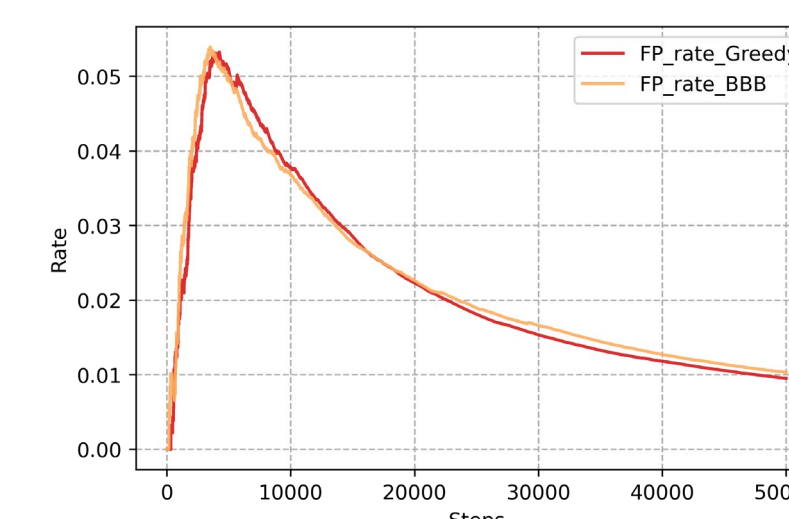


Figure 5. Comparison of false positive rate for greedy and BBB agents

## Observations

- Significantly slower training times than SGD; LRT helps
- Do not observe the original paper's weight distributions

## Next Steps

|                      |                 |
|----------------------|-----------------|
| Weight Pruning       | Laplacian Prior |
| Adversarial Examples |                 |

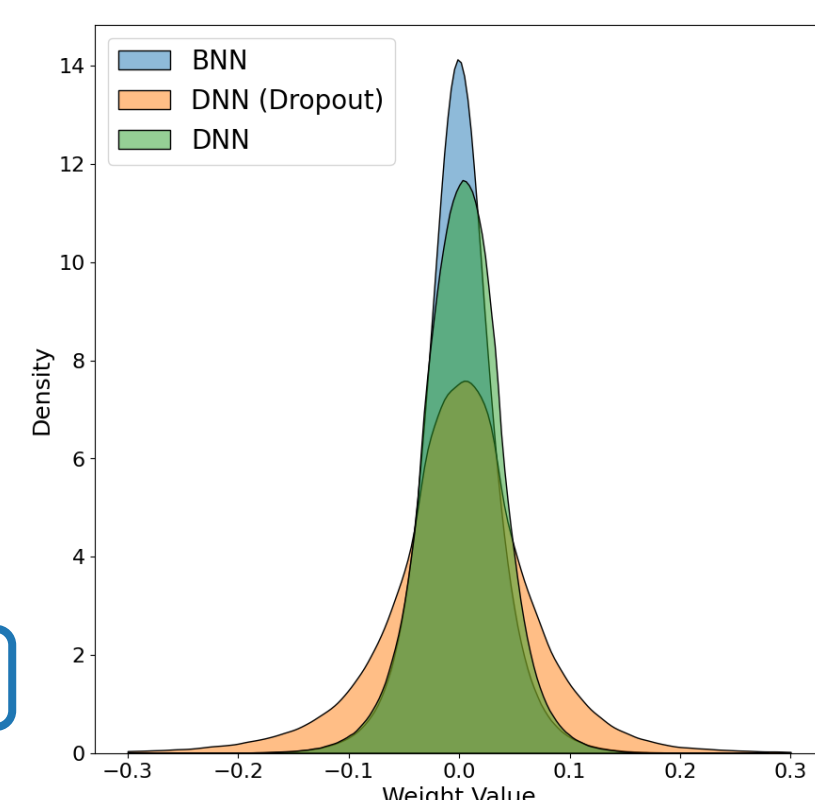


Figure 6. Weight distributions

## References

- [1] Charles Blundell, Julien Cornebise, Koray Kavukcuoglu, and Daan Wierstra, "Weight Uncertainty in Neural Networks," 32nd Int. Conf. Mach. Learn. ICML 2015, vol. 2, pp. 1613–1622, May 2015
- [2] Diederik P. Kingma, Tim Salimans, and Max Welling, "Variational Dropout and the Local Reparameterization Trick," Adv. Neural Inf. Process. Syst., vol. 2015-January, pp. 2575–2583, Jun. 2015

