



45722



Library
Brainware University
398, Ramkrishnapur Road, Barasat
Kolkata, West Bengal-700125

BRAINWARE UNIVERSITY

Term End Examination 2025-2026

Programme – M.Tech.(CSE)-AIML-2025

Course Name – Neural Networks and Deep Learning

Course Code - MTA20502

(Semester II)

Full Marks : 60

Time : 2:30 Hours

[The figure in the margin indicates full marks. Candidates are required to give their answers in their own words as far as practicable.]

Group-A

(Multiple Choice Type Question)

1 x 15=15

1. Choose the correct alternative from the following :

- (i) Learning in ANN occurs through adjustment of
 - a) Activation functions
 - b) Weights and biases
 - c) Number of layers
 - d) Data source
- (ii) Training set consists of
 - a) Input only
 - b) Target only
 - c) Input-output pairs
 - d) Noise
- (iii) BAM stability occurs when energy
 - a) Increases
 - b) Decreases to minimum
 - c) Constant
 - d) Fluctuates
- (iv) In competitive learning, neurons compete based on
 - a) Output error
 - b) Activation strength
 - c) Learning rate
 - d) Bias value
- (v) Counter Propagation Network combines
 - a) Hebbian and delta learning
 - b) Supervised and unsupervised learning
 - c) Recurrent and feedforward
 - d) Clustering and regression
- (vi) Special-purpose neural networks are designed for
 - a) General tasks
 - b) Specific applications
 - c) Random learning
 - d) Universal approximation
- (vii) Hidden units enable neural networks to
 - a) Store datasets
 - b) Learn non-linear relationships
 - c) Reduce parameters
 - d) Eliminate loss functions
- (viii) Reverse-mode automatic differentiation is especially efficient for
 - a) Few inputs and many outputs
 - b) Many inputs and one output
 - c) Linear models only
 - d) Shallow networks
- (ix) Norm penalties can be interpreted as

- a) Hard constraints on outputs
- b) Constrained optimization problems
- c) Probabilistic loss functions
- d) Data preprocessing techniques
- (x) An under-constrained learning problem typically has
 - a) Too much data
 - b) Too many constraints
 - c) More parameters than data
 - d) No hidden layers
- (xi) Data augmentation is especially useful in
 - a) Linear regression
 - b) Image and speech tasks
 - c) Convex optimization
 - d) Feature selection
- (xii) Sparse representations mean
 - a) Most parameters are large
 - b) Most parameters are zero or near zero
 - c) All features are used equally
 - d) Dense activations
- (xiii) Exploding gradients lead to
 - a) Very slow learning
 - b) Weight values growing uncontrollably
 - c) Zero gradients
 - d) Perfect convergence
- (xiv) He initialization is especially suitable for networks using
 - a) Sigmoid activation
 - b) Tanh activation
 - c) ReLU activation
 - d) Linear activation
- (xv) Bayesian optimization is preferred because it
 - a) Requires no evaluations
 - b) Randomly searches space
 - c) Uses probabilistic models of performance
 - d) Ignores past results

Group-B

(Short Answer Type Questions)

3 x 5=15

- 2. Explain Adam optimizer. (3)
- 3. Describe the role of the neighborhood function in the SOFM training algorithm. (3)
- 4. Explain the mechanism of L2 regularization. (3)
- 5. Explain the effect of the value of learning rate in training of neural network. (3)
- 6. Distinguish in between McCulloch-Pitts and Adaline. (3)

OR

- Distinguish in between McCulloch-Pitts and Perceptron. (3)

Group-C

(Long Answer Type Questions)

5 x 6=30

- 7. Explain the concept of Boltzman machine. (5)
- 8. Summarize the structure of a counter propagation network (CPN). (5)
- 9. Explain the method of improving model robustness and regularization by adding noise to the input data during training. (5)
- 10. Discuss Autoencoders and its utility in deep learning. (5)
- 11. Explain the architecture of Madaline. (5)
- 12. Distinguish Local Minima and Saddle Points. (5)

OR

- Differentiate between Xavier (Glorot) and HE Initialization. (5)


Brainware University
 398, Ramkrishnapur Road, Barasat
 Kolkata, West Bengal-700125