

Code No: 156BN

R18

JAWAHARLAL NEHRU TECHNOLOGICAL UNIVERSITY HYDERABAD

B. Tech III Year II Semester Examinations, August/September - 2021

MACHINE LEARNING

(Computer Science and Engineering)

Time: 3 Hours

Max. Marks: 75

Answer any five questions

All questions carry equal marks

1. a) Define Well-Posed problem. Illustrate any four examples for well-posed problems.
b) What do you mean by Candidate elimination? Explain. [7+8]
2. a) What are the concepts of learning as search? Discuss.
b) Discuss the appropriate problems for decision tree learning. [8+7]
3. a) Contrast the hypothesis space search in ID3 and the candidate elimination algorithm. [7+8]
b) Explain the Backpropagation learning algorithm and its limitations.
4. a) How a multi-layered network learn using a gradient descent algorithm? Discuss. [8+7]
b) Explain the methods for comparing the accuracy of two hypotheses.
5. a) State Bayes theorem. Illustrate Bayes theorem with an example.
b) Describe the mistake-bound model of learning. [8+7]
6. a) Explain Gibbs algorithm with an example.
b) State and explain the Minimum Description Length Principle. [8+7]
7. a) Discuss about Hypothesis space search in genetic algorithms.
b) Write the basic algorithm for learning sets of First-Order Rules. [8+7]
8. a) Discuss Explanation-Based learning of search control knowledge.
b) Explain the inductive analytical approaches to learning. [8+7]

---ooOoo---

ANSWER KEY

1. a) Define a Well-Posed problem. Illustrate any four examples for Well-Posed problems.

1. a) Define Well-Posed Problem

1. Definition: A well-posed problem in the context of computational theory or mathematical analysis is one that meets three criteria: it has a solution, the solution is unique, and the solution's behavior changes continuously with the initial conditions.

2. Existence of Solution: The problem must have at least one solution.

3. Uniqueness of Solution: There should be no ambiguity; only one solution must exist for the given problem.

4. Stability: The solution must not change drastically with small changes in the initial conditions.

5. Real-world Applicability: Well-posed problems are crucial in real-world scenarios for predictable and reliable outcomes.

6. Example 1 - Linear Equations: Solving a system of linear equations is a well-posed problem as it typically has a unique solution.

7. Example 2 - Classical Physics Problems: Problems in classical mechanics, where initial conditions can determine the motion of a particle uniquely and continuously.

8. Example 3 - Calculating Compound Interest: Given the principal, rate, and time, the calculation of compound interest is well-posed.

9. Example 4 - Finding Roots of a Polynomial: Determining the roots of a non-degenerate polynomial equation.

10. Practical Significance: In AI and computational problems, ensuring a problem is well-posed is fundamental for designing algorithms that provide reliable solutions.

1. b) What do you mean by Candidate elimination? Explain.

1. Definition: Candidate elimination is a method used in machine learning to identify a hypothesis that correctly predicts the classification of instances in a given dataset.

2. Concept of Version Space: It involves maintaining the version space, which is the set of all hypotheses consistent with the training examples.

3. General-to-Specific Ordering: The algorithm iterates through the hypothesis space, starting from the most general hypothesis and moving towards more specific ones.

4. Process: It systematically eliminates hypotheses that fail to classify the training examples correctly.

5. Use in Concept Learning: Primarily used in concept learning to find the optimal hypothesis that can be generalized from examples.

6. Handling Noise: It is robust in handling noise and errors in the training data.

7. Advantage: Helps in understanding the boundaries of what can be learned from the given data.
8. Algorithm Complexity: The complexity can vary based on the size of the hypothesis space and the number of attributes.
9. Practical Applications: Used in various machine learning tasks like pattern recognition, spam detection, and decision-making processes.
10. Example: In spam email detection, candidate elimination helps to refine the hypothesis about what constitutes spam versus non-spam emails.

2. a) What are the concepts of learning as search? Discuss.

1. Definition: Learning as search is a concept where the learning process is viewed as a search through a space of possible hypotheses or solutions to find the one that best fits the data.
2. Hypothesis Space: Involves defining a space of possible hypotheses that could explain the observed data. The learning algorithm searches through this space.
3. Search Algorithms: Various algorithms like hill climbing, genetic algorithms, or exhaustive search can be used to navigate the hypothesis space.
4. Optimization Goal: The goal is to find the hypothesis that maximizes (or minimizes) a particular objective function, such as accuracy or error rate.
5. Overfitting and Underfitting: Balancing between overfitting (too complex hypotheses fitting noise in the data) and underfitting (too simple hypotheses missing important patterns).
6. Version Space: Subset of the hypothesis space that is consistent with the training examples. The search process often involves narrowing down this version space.
7. Pruning Strategies: Techniques like pruning are used to reduce the search space and avoid overfitting.
8. Incremental Learning: Some learning-as-search approaches update hypotheses incrementally as new data arrives, rather than processing all data in batch mode.
9. Complexity and Efficiency: The efficiency of the learning process depends on the size of the hypothesis space and the efficiency of the search algorithm.
10. Example-Based Learning: Approaches like nearest neighbor methods can be viewed as a search for the most similar past examples to a new instance.

2. b) Discuss the appropriate problems for decision tree learning.

1. Classification Problems: Decision trees are well-suited for classification tasks where the objective is to assign each input to one of the predefined categories.
2. Handling Discrete and Continuous Features: They can handle both types of features efficiently.
 1. Non-Linear Relationships: Good at capturing non-linear relationships between features and the target variable.

4. Missing Values: Capable of handling missing values in the data.
5. Interpretability: Provide clear and interpretable models, making them useful in applications where understanding the model's decision process is important.
6. Large Datasets: Suitable for large datasets, though the size of the tree can become a constraint.
7. Preprocessing: Minimal data preprocessing is needed compared to other algorithms.
8. Feature Selection: Automatically perform feature selection by choosing the most informative features for splits.
9. Multi-Output Problems: Can be extended to handle multi-output problems (predicting multiple dependent variables).
10. Non-parametric Method: As a non-parametric method, they don't require assumptions about the distribution of the data.

3.a) Contrast the hypothesis space search in ID3 and candidate elimination algorithm.

1. Algorithm Type: ID3 is a decision tree learning algorithm, whereas the Candidate Elimination algorithm is a version space learning algorithm.
2. Hypothesis Space Representation: ID3 represents the hypothesis space through a decision tree, while Candidate Elimination represents it as a set of consistent hypotheses bounded by general and specific boundaries.
3. Handling Noise: ID3 is less robust to noisy data compared to the Candidate Elimination algorithm which is more tolerant of noise.
4. Data Requirements: ID3 requires a complete dataset for effective training. Candidate Elimination can handle partially classified datasets better.
5. Learning Process: ID3 incrementally builds a decision tree by selecting the best attribute using information gain. Candidate Elimination updates the general and specific boundaries of the hypothesis space as it receives new examples.
6. Output: The output of ID3 is a single decision tree. In contrast, Candidate Elimination outputs a version space that includes all hypotheses consistent with the examples.
7. Pruning Strategy: ID3 may require post-pruning to handle overfitting. Candidate Elimination inherently avoids overfitting by maintaining the version space.
8. Complexity and Efficiency: ID3 is generally faster and less computationally complex than Candidate Elimination, which can become computationally expensive with a large hypothesis space.
9. Assumptions: ID3 assumes attributes are categorical and discretizes continuous attributes. Candidate Elimination doesn't require this discretization.
10. Usage Context: ID3 is preferred for problems where a single model is sufficient, whereas Candidate Elimination is suitable for situations where maintaining multiple consistent hypotheses is beneficial.

3. b) Explain the Backpropagation Learning Algorithm and Its Limitations

1. Definition: Backpropagation is a supervised learning algorithm used for training artificial neural networks. It involves adjusting the weights of the network to minimize the error between the actual output and the desired output.
2. Algorithm Process: It uses gradient descent to update network weights iteratively based on the error rate obtained in the previous epoch (cycle).
3. Learning Rate: The learning rate is a crucial parameter that affects the convergence of the algorithm. An inappropriate rate can lead to slow convergence or overshooting the minimum.
4. Problem of Local Minima: Backpropagation can get stuck in local minima, leading to suboptimal solutions.
5. Vanishing Gradient Problem: In deep networks, the gradient can become vanishingly small, preventing the weights in the initial layers from changing significantly.
6. Overfitting: Without proper regularization, backpropagation can lead to overfitting, where the model performs well on training data but poorly on unseen data.
7. Computational Intensity: The algorithm can be computationally intensive, especially for deep networks with many layers and neurons.
8. Data Requirements: Requires a large amount of training data to perform effectively and avoid overfitting.
9. Initialization Sensitivity: The initial choice of weights can significantly affect the final converged solution.
10. Feature Scaling Necessity: Input features need to be normalized or standardized for the algorithm to work effectively.

4.a) How a multi layered network learns using a gradient descent algorithm? Discuss.

1. Understanding the Basics: A multi-layered network, commonly known as a neural network, consists of layers of interconnected nodes or neurons. Each neuron performs a simple calculation.
2. Input Processing: The network takes input data and processes it through multiple layers, transforming the input at each step based on the current weights of the connections.
3. Forward Propagation: During forward propagation, data moves from the input layer through the hidden layers to the output layer, producing a prediction.
4. Error Calculation: The network calculates the error by comparing the prediction with the actual target values, often using a loss function like Mean Squared Error.
5. Backpropagation: Gradient descent comes into play during backpropagation. The algorithm calculates the gradient of the loss function with respect to each weight by the chain rule.

6. **Gradient Calculation:** The gradient indicates the direction in which the weights need to be adjusted to minimize the error.
7. **Weight Update:** Weights are updated in the opposite direction of the gradient, scaled by a learning rate, to gradually reduce the error.
8. **Iterative Process:** This process of forward propagation, error calculation, backpropagation, and weight update is repeated for many iterations or epochs.
9. **Learning Rate Importance:** The learning rate is a crucial hyperparameter that determines the size of the steps taken in the weight space. A too high rate can overshoot the minimum, while a too low rate can slow down the training process.
10. **Convergence:** The process continues until the network's performance no longer significantly improves, indicating that it has learned to map inputs to outputs effectively.
11. **Regularization and Optimization:** Techniques like regularization (e.g., L1, L2) and advanced optimizers (e.g., Adam, RMSprop) are often used to enhance the learning process and prevent overfitting.
12. **Generalization Ability:** The ultimate goal is for the network to not just memorize the training data but to generalize well to new, unseen data.
13. **Visualization and Analysis:** Tools and techniques like learning curves and gradient visualization help in understanding and improving the learning process.
14. **Practical Applications:** This learning mechanism is behind many modern AI applications, from image recognition to natural language processing.
15. **Continuous Evolution:** The field of neural networks and gradient descent is continuously evolving, with new research leading to more efficient and effective learning methods.

4.b) Explain the Methods for Comparing the Accuracy of Two Hypotheses

1. **Confusion Matrix:** Use a confusion matrix for each hypothesis to visualize true positives, true negatives, false positives, and false negatives.
2. **Accuracy Score:** Calculate the accuracy as the ratio of correctly predicted instances to the total instances. Compare these scores directly.
3. **Precision and Recall:** Evaluate precision ($\text{true positives} / (\text{true positives} + \text{false positives})$) and recall ($\text{true positives} / (\text{true positives} + \text{false negatives})$) for a balanced assessment of performance, especially in imbalanced datasets.
4. **F1 Score:** The F1 Score combines precision and recall into a single metric, which is the harmonic mean of precision and recall, providing a balance between the two.
5. **ROC-AUC Curve:** The Receiver Operating Characteristic (ROC) curve plots the true positive rate against the false positive rate. The Area Under the Curve (AUC) provides a single value metric to compare models.
6. **Cross-Validation:** Use cross-validation to assess how the hypotheses perform on different subsets of the dataset, ensuring robustness and reducing the likelihood of overfitting.

7. **Statistical Tests:** Apply statistical hypothesis testing, like t-tests or ANOVA, to determine if the difference in performance is statistically significant.
 8. **Learning Curves:** Analyze learning curves for training and validation data to understand how well each hypothesis generalizes.
 9. **Error Analysis:** Perform a qualitative error analysis to understand the types of errors each hypothesis makes, which might be crucial depending on the application.
 10. **Practical Implications:** Beyond accuracy, consider factors like computational efficiency, scalability, and ease of interpretation, which can be crucial in practical applications.
 11. **Domain-Specific Metrics:** In some fields, domain-specific metrics might be more relevant than generic accuracy measures. For example, in medical diagnosis, metrics like sensitivity and specificity are often more important.
 12. **User Feedback:** Sometimes, user feedback or expert evaluation can provide insights into the model's performance, especially in subjective tasks like content generation.
 13. **Ensemble Learning:** Investigate if combining the hypotheses through ensemble methods like stacking, boosting, or bagging improves overall performance.
 14. **Bias-Variance Tradeoff:** Analyze each hypothesis in terms of bias (error from erroneous assumptions) and variance (error from sensitivity to small fluctuations in the training set).
 15. **Dataset Sensitivity:** Evaluate how sensitive each hypothesis is to changes in the dataset, which can reveal dependencies and robustness issues.
- Each method offers unique insights, and often, a combination of these methods provides a comprehensive understanding of the comparative performance of two hypotheses.

5. a) State Bayes theorem. Illustrate Bayes theorem with an example.

Bayes Theorem Statement:

1. **Fundamental Concept:** Bayes theorem describes the probability of an event, based on prior knowledge of conditions that might be related to the event.
2. **Mathematical Formulation:** It is mathematically stated as $P(A|B) = \frac{P(B|A) \times P(A)}{P(B)}$.
3. **Components Explained:**
 - $P(A|B)$ is the probability of event A occurring given that B is true.
 - $P(B|A)$ is the probability of event B given that A is true.
 - $P(A)$ and $P(B)$ are the probabilities of observing A and B independently of each other.

Example:

Imagine a medical context where:

Event A: Having a certain disease.

Event B: Testing positive for that disease.

1. Known Probabilities:

The probability of having the disease ($P(A)$) is 0.01 (1%).

The probability of testing positive if you have the disease ($P(B|A)$) is 0.95 (95%).

- The probability of testing positive in general ($P(B)$) is 0.02 (2%).

2. Applying Bayes Theorem:

We calculate $P(A|B)$, the probability of having the disease given a positive test result.

Use the formula: $P(A|B) = \frac{P(B|A) \times P(A)}{P(B)} = \frac{0.95 \times 0.01}{0.02} = 0.475$.

3. Interpretation:

Despite a high test accuracy (95%), the probability of actually having the disease after testing positive is only 47.5%.

5. b) Describe the mistake-bound model of learning.

Concept of Mistake Bound Model:

1. Learning Framework: The mistake-bound model is used in machine learning to evaluate the performance of online learning algorithms.

2. Mistake-Driven Learning: It focuses on how many mistakes an algorithm makes before it learns the correct hypothesis.

3. Mistake Bound: This is the upper limit on the number of mistakes the algorithm is allowed to make.

Key Points:

4. Online Learning: Algorithms learn by sequentially processing instances and making predictions.

5. Immediate Feedback: After each prediction, the algorithm receives immediate feedback about its accuracy.

6. Update Strategy: Based on feedback, the algorithm updates its model/hypothesis.

7. Performance Measure: The quality of an algorithm is measured by the number of errors (mistakes) it makes before converging to the correct hypothesis.

8. Example Algorithms: The Perceptron and Winnow algorithms are classic examples of learning with a mistake bound.

9. Applicability: Particularly useful in scenarios where data arrives in a stream, and immediate adaptation is necessary.

10. Limitation: Assumes that there exists a hypothesis within the algorithm's hypothesis class that can make the correct prediction for all instances.

In this model, the goal is to minimize the mistake bound, which reflects the efficiency and effectiveness of the learning process.

6. a) Explain Gibbs algorithm with an example.

1. **Introduction to Gibbs Algorithm:** Gibbs algorithm, also known as Gibbs sampling, is a Markov Chain Monte Carlo (MCMC) method used for statistical inference and generating samples from complex probability distributions.
2. **Basic Principle:** It works by sequentially updating each variable in a set, conditioning on the current values of all other variables. This process creates a Markov chain of samples.
3. **Algorithm Steps:**
 - Initialize all the variables with some values.
 - For each iteration:
 - Select a variable to update.
 - Keep all other variables fixed.
 - Sample a new value for the selected variable from its conditional distribution.
4. **Convergence:** Over many iterations, the distribution of the variables converges to the joint distribution of the system.
5. **Example Scenario:** Consider a probabilistic model with two variables, X and Y, each taking binary values (0 or 1). The joint probability distribution is known.
6. **Initialization:** Start with arbitrary values, say $X=0$ and $Y=0$.
7. **Iteration Process:**
 - Update X: Sample a new value for X from $P(X|Y=0)$.
 - Update Y: Sample a new value for Y from $P(Y|X=\text{new value of } X)$.
8. **Repetition:** Repeat the above steps multiple times. Each complete set of updates constitutes one iteration.
9. **Convergence Check:** After several iterations, the samples of X and Y will represent samples from the joint distribution $P(X, Y)$.
10. ****Applications**:** Gibbs sampling is widely used in Bayesian statistics, machine learning, and in scenarios where direct sampling is difficult or impossible.

6.b) State and explain the Minimum Description Length Principle.

1. **Introduction:** The Minimum Description Length (MDL) Principle is a formalization of Occam's Razor in information theory and statistics.
2. **Basic Concept:** MDL advocates for choosing the model that gives the shortest overall description of the data and the model itself.
3. **Two Components of MDL:**
 - Length of the description of the model (complexity of the model).
 - Length of the description of the data when encoded with the model (goodness of fit).
4. **Trade-off:** MDL balances model complexity against the fit of the model to the data. A complex model might fit the data better but requires a longer description.

5. **Model Selection:** In MDL, the best model minimizes the sum of the length of encoding the model and the length of encoding the data given the model.
6. **Encoding:** It refers to how both the model and the data are represented. The idea is akin to compressing the data with the model.
7. **Example:** Consider fitting a polynomial to data points. A higher-degree polynomial might reduce error but is more complex. MDL helps to decide the degree of the polynomial by balancing complexity and fit.
8. **Practical Application:** MDL is widely used in machine learning and data compression, guiding the choice of models or algorithms that generalize well from training data to unseen data.
9. **Advantage:** It provides a principled approach to avoid overfitting, especially in situations with limited data.
10. **Philosophical Foundation:** The MDL Principle is rooted in the idea that simpler models are more likely to be true, reflecting a preference for simplicity in scientific theories.

7.a) Discuss about Hypothesis Space Search in Genetic Algorithms.

1. **Definition:** Hypothesis space search in genetic algorithms (GAs) involves exploring a set of potential solutions (hypotheses) to optimize a particular problem.
2. **Genetic Representation:** Solutions are represented as chromosomes (strings), often binary, mimicking biological DNA.
3. **Initial Population:** The search begins with a randomly generated population of hypotheses.
4. **Fitness Function:** Each hypothesis is evaluated using a fitness function, determining how close it is to the optimal solution.
5. **Selection:** Hypotheses with higher fitness are more likely to be selected for reproduction.
6. **Crossover:** Selected hypotheses undergo crossover, where segments of their strings are exchanged, to produce new hypotheses.
7. **Mutation:** Some parts of the new hypotheses may randomly mutate, introducing variability.
8. **Replacement:** New hypotheses replace some or all of the old population, based on their fitness.
9. **Iteration:** This process repeats over several generations, with each iteration potentially improving the hypotheses.
10. **Convergence:** The algorithm converges when an optimal or satisfactory hypothesis is found or after a set number of generations.

7.b) Write the Basic Algorithm for Learning Sets of First-Order Rules.

1. **Initialization:** Start with an empty set of rules.
2. **Select Target Concept:** Choose a specific concept to learn about.

3. Positive and Negative Examples: Collect examples that either satisfy (positive) or do not satisfy (negative) the target concept.
4. Rule Generation: Begin generating rules that explain the positive examples while not covering negative ones.
5. Hypothesis Space: Consider the space of all possible rules that can be formed using the attributes of the examples.
6. Specialization: Start with a general rule and iteratively add conditions to specialize it, ensuring it fits the positive examples better.
7. Evaluation: Evaluate the rules based on their accuracy in classifying the examples.
8. Pruning: Remove unnecessary conditions from rules to make them simpler and more general.
9. Iterative Refinement: Repeatedly refine the rules by comparing against examples and adjusting accordingly.
10. Termination: The algorithm terminates when it finds a set of rules that adequately covers all positive examples and excludes negative ones.

8.a) Explanation-Based Learning (EBL) of Search Control Knowledge.

1. Definition: Explanation-Based Learning focuses on improving search control in problem-solving by learning from specific examples.
2. Goal-oriented: EBL aims to extract general rules from specific problem-solving instances to improve future search efficiency.
3. Knowledge Intensive: Relies heavily on existing domain knowledge to analyze new examples.
4. Explanation Construction: Involves building an explanation for why a particular solution works for a given problem.
5. Generalization: The specific explanation is then generalized to apply to a broader set of problems.
6. Reduction of Search Space: By applying learned knowledge, EBL reduces unnecessary exploration in future problem-solving.
7. Focus on Relevance: EBL emphasizes learning only relevant information, improving the effectiveness of the learned knowledge.
8. Integration with Other Learning Methods: Often used in combination with other learning techniques for more robust learning.
9. Application Examples: Used in complex domains like medical diagnosis, chess playing, and scientific discovery.
10. Challenges: Requires substantial domain knowledge and can be limited by the quality of explanations.

8. b) Inductive Analytical Approaches to Learning.

1. Definition: Inductive analytical learning involves inferring general rules from specific instances, using analytical reasoning.

2. Data-Driven Learning: Starts with examples and infers general rules or patterns.
3. Hypothesis Formation: Involves creating hypotheses about the underlying rules governing the examples.
4. Analytical Reasoning: Uses logical reasoning to validate or refine these hypotheses.
5. Incorporation of Background Knowledge: Often integrates existing domain knowledge to guide hypothesis formation.
6. Iterative Refinement: Involves iteratively testing and refining hypotheses based on new data.
7. Handling Noise and Exceptions: Capable of dealing with noisy data and exceptions through robust hypothesis testing.
8. Application Areas: Widely used in machine learning, natural language processing, and data mining.
9. Empirical Evaluation: Relies on empirical methods to evaluate the correctness and applicability of the learned rules.
10. Challenges: Requires a balance between overfitting (too specific) and underfitting (too general) of the learned rules.