

인공지능

1. 다음 설명에 해당하는 기계 학습의 유형은?

어떤 환경 안에서 정의된 에이전트가 현재의 상태를 인식하여, 선택 가능한 행동들 중 보상을 최대화하는 행동 혹은 행동 순서를 선택하는 방법

- ① 강화 학습(reinforcement learning)
- ② 지도 학습(supervised learning)
- ③ 비지도 학습(unsupervised learning)
- ④ 앙상블 학습(ensemble learning)

2. 이진 분류기의 성능평가에 대한 설명으로 옳은 것은?

- ① 분류기가 양성으로 판정한 것들 중에서 실제 양성인 것의 비율을 재현율(recall)이라고 한다.
- ② 실제 음성인 것들 중에서 분류기가 음성으로 판정한 것의 비율을 특이도(specificity)라고 한다.
- ③ 실제 양성인 것들 중에서 분류기가 양성으로 판정한 것의 비율을 위양성율(false positive rate)이라고 한다.
- ④ 분류기는 민감도(sensitivity) 값이 작을수록 좋고 위양성율 값이 클수록 좋다.

3. 운전면허 시험 성적에서 합격 6명, 불합격 4명으로 구성된 데이터 세트의 지니 계수(Gini index) 값은?

- ① 0.25
- ② 0.48
- ③ 0.52
- ④ 1

4. 선형 회귀에서 평균제곱오차(MSE)를 최소화하는 계수를 찾는 방법은?

- ① 주성분 분석법
- ② 선형 판별 분석법
- ③ 최소우도법
- ④ 최소제곱법

5. 시간성을 갖는 데이터에 유용하고 문맥 의존성을 효율적으로 처리할 수 있는 신경망은?

- ① CNN
- ② KNN
- ③ MLP
- ④ RNN

6. 탐색 방법에 대한 설명으로 옳지 않은 것은?

- ① 맹목적 탐색은 문제의 상태 공간 정보를 이용한다.
- ② 깊이 우선 탐색은 너비 우선 탐색에 비해 메모리에 대한 부담이 적지만 최단 경로를 찾는다는 보장이 없다.
- ③ 너비 우선 탐색은 목표 상태를 찾을 때까지 생성된 모든 노드를 메모리에서 관리하기 때문에 메모리 비용이 크다.
- ④ 반복적 깊이 심화 탐색은 탐색 깊이 한계를 0부터 점차 증가 시키면서 깊이 우선으로 탐색하는 방법이고 최단 경로 찾는 것을 보장한다.

7. 추천 시스템(recommender system)의 유형 중 협업 기반 추천과 내용 기반 추천에 대한 설명으로 옳지 않은 것은?

- ① 내용 기반 추천은 메모리 기반 방법과 모델 기반 방법이 있고 콜드 스타트(cold start) 문제가 발생한다.
- ② 아이템 기반 추천은 이전에 구매한 아이템을 기반으로 그 상품과 유사한 다른 상품을 추천하는 방법이다.
- ③ 사용자 기반 추천은 나와 비슷한 성향을 지닌 사용자의 데이터를 기반으로 상품을 추천하는 방법이다.
- ④ 추천 시스템은 콘텐츠의 내용에 기반하거나 사람들의 성향 정보를 취득한 후 개인화된 맞춤 항목을 추천한다.

8. 규칙 기반 시스템에 대한 설명으로 옳지 않은 것은?

- ① 경합 해소(conflict resolution)는 여러 규칙을 동시에 실행할 경우 모순이 발생할 수 있기 때문에 실행 가능한 규칙들 중에서 한 개를 선택하는 것을 의미한다.
- ② 전향 추론(forward chaining)은 규칙의 조건부를 만족시키는 사실이 있을 때 규칙의 결론부를 실행하거나 처리한다.
- ③ 후향 추론(backward chaining)은 특정 사실을 확인하기 위해 해당 사실을 결론부에 포함하는 규칙의 조건부가 만족하는지 확인해 가는 방식으로 추론한다.
- ④ 추론 엔진(inference engine)은 경합 해소 → 패턴 매칭(pattern matching) → 규칙 실행(rule execution)의 과정을 반복하여 추론한다.

9. A* 알고리즘에 대한 설명으로 옳지 않은 것은?

- ① A* 알고리즘의 평가함수는 $f(n) = g(n) + h(n)$ 으로 표현한다.
- ② 휴리스틱 함수 $\hat{h}(n)$ 이 실제 남은 비용을 과대평가하지 않으면 허용성(admissibility)을 만족한다.
- ③ 시작 노드에서 노드 n 까지의 실제 비용 $g(n)$ 은 탐색이 진행됨에 따라 감소할 수 있다.
- ④ 휴리스틱 함수 $\hat{h}(n)$ 이 허용성을 만족하면 A* 알고리즘은 항상 최적해를 찾는다.

10. 이진 분류 문제에 대한 혼동행렬(confusion matrix)에서 각각 계산한 민감도, 정확도(accuracy), 위양성율 값을 모두 곱한 값은?

		예측값	
		양성	음성
실제값	양성	40	10
	음성	15	35

- $$\begin{array}{ll} \textcircled{1} \quad \frac{3}{80} & \textcircled{2} \quad \frac{32}{275} \\ \textcircled{3} \quad \frac{3}{55} & \textcircled{4} \quad \frac{9}{50} \end{array}$$

11. 기계 학습에 대한 설명으로 옳지 않은 것은?

- ① 군집화와 밀도추정 문제는 일반적으로 비지도 학습 방식을 사용한다.
- ② 지도 학습의 학습 데이터는 입력과 기대 출력 정보의 쌍으로 구성된다.
- ③ 데이터 증강은 입력 데이터를 저차원 데이터로 압축하는 인코더 형태의 모델을 통해 학습한다.
- ④ 프로그램을 명시적으로 작성하지 않고 컴퓨터가 학습할 수 있는 능력을 갖추게 하는 기술이다.

12. 생성적 적대 신경망(generative adversarial networks)에서 생성자(generator)와 판별자(discriminator)에 대한 설명으로 옳은 것만을 모두 고르면?

- ㄱ. 생성자와 판별자는 적대적인 관계에서 경쟁하면서 훈련한다.
- ㄴ. 생성자는 진짜 데이터를 만들면서 판별자가 가짜라고 속을 때까지 학습한다.
- ㄷ. 판별자는 훈련 데이터를 진짜 데이터로 판별하고 생성자가 만든 데이터는 가짜 데이터로 판별하도록 학습한다.

- ① $\neg$, $\perp$
② $\neg$, $\sqsubset$
③ $\perp$, $\sqsubset$
④ $\neg$, $\perp$, $\sqsubset$

13. 역학조사에서 흡연 여부와 암 유병 여부를 조사하였더니 다음 표와 같이 계산되었다. 흡연하는 사람일 사건을 A로 나타내고 암환자일 사건을 B로 나타내었을 때, 결합 확률 $P(A \cap B)$ 의 값은?

	비흡연	흡연	합계(명)
정상인	35	15	50
암환자	7	3	10
합계(명)	42	18	60

- ① 0.05
- ② 0.15
- ③ 0.25
- ④ 0.33

14. 빅데이터 분석을 위한 데이터 마이닝에 대한 설명으로 옳지 않은 것은?

- ① 정제되지 않은 대용량의 데이터로부터 의미 있는 상관관계, 패턴, 추세 등을 발견한다.
- ② 트랜잭션에서 항목 간의 불순도를 표현하기 위해 최소지지도와 최소신뢰도를 사용한다.
- ③ 저장된 데이터를 정제(cleaning), 통합(integration)하여 잡음과 불일치를 제거하고 데이터 웨어하우스에 저장한다.
- ④ 대규모 데이터에서 암묵적인, 이전에 알려지지 않은, 잠재적으로 유용할 것 같은 정보나 지식을 추출하는 체계적인 과정이다.

15. 다음 설명의 (가)에 들어갈 학습 방법은?

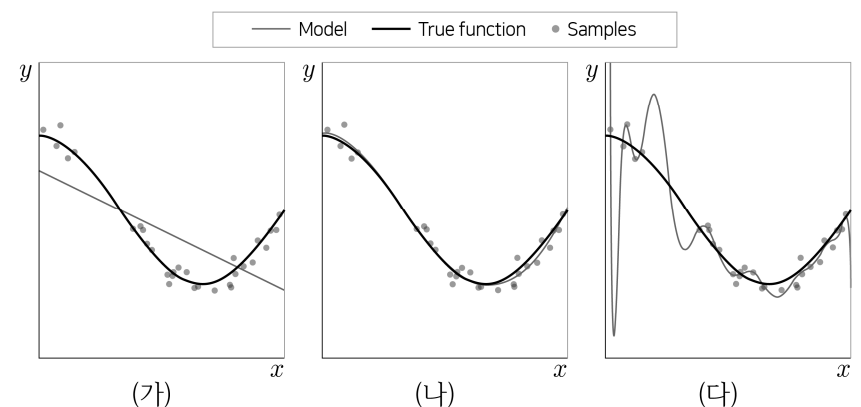
(가)은 학습하는 방법을 학습하는 것을 의미한다. 특정 작업을 인간과 같이 빠르게 학습하는 능력을 만들기 위해 동일한 분포를 이루는 여러 작업에 대한 학습 방법을 학습한 뒤, 유사한 작업을 적은 데이터로 빠르고 효과적으로 학습할 수 있게 한다.

- ① 연합 학습
 - ② 메타 학습
 - ③ 강화 학습
 - ④ 표현 학습

16. 결정 트리 학습 알고리즘에 대한 설명으로 옳지 않은 것은?

- ① C4.5, C5.0, CHAID, CART 등의 알고리즘이 있다.
- ② 엔트로피는 어떤 집단에 각 부류의 데이터들이 균등한 비율로 섞여 있을수록 작아지는 특성이 있다.
- ③ ID3 알고리즘은 범주형 속성값을 갖는 학습 데이터로부터 엔트로피 개념을 사용하여 결정 트리를 만든다.
- ④ 정보 이득(information gain)은 특정한 속성이 원하는 분류 방식에 부합하게 데이터를 나누는지를 측정할 수 있는 척도이다.

17. 다음 그래프에 대한 설명으로 옳지 않은 것은? (단, x 는 특징, y 는 레이블이다)



- ① (가)는 모델의 표현력이 부족해 훈련 오차와 검증 오차가 모두 높게 유지되는 상황을 나타낸다.
- ② (나)는 훈련 오차와 검증 오차가 모두 낮고 서로 차이가 작아, 모델이 데이터에 일반화되고 있음을 보여 준다.
- ③ (다)는 훈련 오차가 낮고 검증 오차는 높게 유지되므로, 성능을 향상하기 위해서 모델 복잡도를 더 높여야 한다.
- ④ 기계 학습에서 과적합은 학습 데이터에서 성능이 뛰어나지만 일반화에 대한 성능을 보장할 수 없다.

