

◆ 10-9평 36~39번

[36~39] 다음 글을 읽고 물음에 답하시오.

매일 쏟아지는 수많은 우편물들은 발송 지역별로 분류되어야 한다. 우편물 분류 작업은 우편번호 숫자를 인식함으로써 자동화될 수 있다. 이때 자동분류기는 환경과의 상호 작용에 기반한 경험적인 데이터로부터 스스로 성능을 향상시킬 수 있는 학습 능력을 갖춰야 한다. ㉠ 학습은 상호 작용의 정도에 따라 경험하는 데이터가 달라지고, 이러한 학습 데이터에 따라 자동분류기의 성능이 달라지게 된다. 즉, 자동분류기는 단순히 데이터를 기억하는 것이 아니라, 다양한 경험에서 새로운 정보를 추론하여 스스로 분류할 수 있는 능력을 갖춰야 한다.

	학습 데이터	실험 데이터
필기체 숫자	5500	5
입력 특징		
목표치	5 5 0 0	

우편번호 자동분류기가 학습하기 위해서는, 먼저 우편번호 숫자를 하나씩 분할하고, 0부터 9까지를 잘 구별할 수 있는 입력 특징을 찾아야 한다. 위 그림은 필기체 숫자를 가로, 세로 8 등분하여 연필이 지나간 자리를 1, 그렇지 않으면 0의 값을 주어, 입력 특징을 추출한 것이다.

다음으로, 추출된 특징으로 학습할 때 분류기에 목표치를 제공함으로써 학습을 감독할 수 있다. 즉, 입력 특징에 대한 목표치가 제시되면 분류기는 데이터를 제시된 목표치로 분류하도록 학습한다. 이렇게 목표치를 이용하는 학습을 ㉡ 감독학습이라 한다. 숫자 분류기에 0부터 9까지 각각의 숫자에 대한 목표치가 제공되면, 분류기는 감독학습을 수행한다. 위의 그림에서 분류기는 네 개의 학습 데이터에 대한 입력 특징과 목표치를 통해 학습한다. 이 학습을 통해 두 개의 '5'와 두 개의 '0'을 각각 같은 숫자로 인식하면서, 동시에 '5'와 '0'을 서로 다른 숫자로 분류해 내는 함수를 만든다. 감독학습을 통해 올바르게 학습하였다면, 그림의 실험 데이터는 숫자 '5'로 인식된다.

그러면, 목표치를 주는 것이 어려운 경우에는 어떻게 학습할까? 목표치가 없을 때는 학습 데이터로 주어진 입력 특징들의 유사성을 찾아 군집화한다. 이와 같이 목표치가 제시되지 않는 학습을 무감독학습이라고 한다. 예컨대 위 그림에서 네 개의 필기체 숫자에 대한 입력 특징만 주어지면, 무감독학습은 비슷한 입력 특징을 가진 숫자들을 ㉢ 모아 '5' 또는 '0'에 대해 군집화하는 함수를 만든다. 무감독학습을 통해 올바르게 학습하였다면, 실험 데이터는 '5'의 군집과 유사한 것으로 인식된다.

이렇게 학습된 자동분류기는 실험 데이터를 정확하게 분류하였는지에 따라 그 성능이 평가된다. 이러한 과정을 통해 우편번호 자동분류기는 우편물을 지역별로 분류할 수 있게 된다.

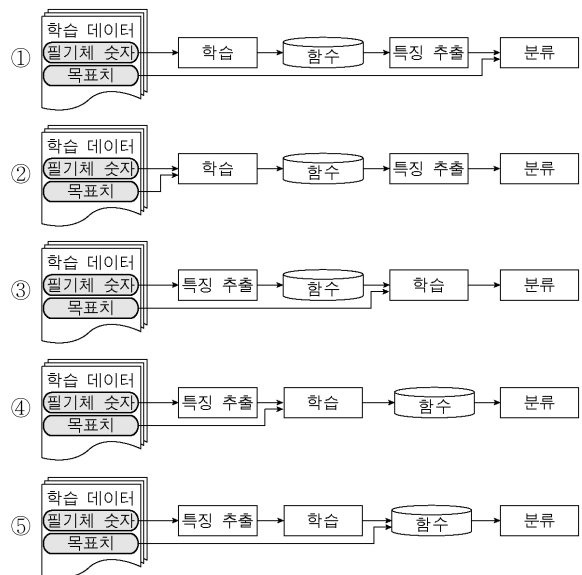
36. 위 글의 '우편번호 자동분류기'에 대한 설명으로 적절한 것은?

- ① 자동분류기의 성능은 학습 데이터의 양에 영향을 받지 않는다.
- ② 우리나라 우편번호 자동분류기는 총 6종류의 목표치를 이용한다.
- ③ 자동분류기의 학습은 일정한 종류의 필기체 숫자를 기억하는 것이다.
- ④ 자동분류기는 0부터 9까지의 차이를 최소화하는 입력 특징을 사용한다.
- ⑤ 자동분류기의 학습은 필기체 숫자의 목표치가 없으면, 유사한 입력 특징을 가진 것끼리 모은다.

37. 휴대 전화의 기능을 소개하는 문구 중, ㉠의 기능을 담은 예로 적절하지 않은 것은? [3점]

- ① 전화가 걸려 오면 등록된 수신 거부 목록과 일일이 대조하여, 목록에 있는 번호이면 수신을 거부한다.
- ② 휴대 전화를 든 손으로 등록된 단축 번호를 공중에 쓰면, 전화기가 숫자를 인식하여 자동으로 전화를 건다.
- ③ 사용자의 음성 특징을 추출하여 사용자와 타인의 음성을 분류하면, 사용자의 음성으로만 휴대 전화를 사용할 수 있다.
- ④ 휴대 전화에 닿는 형태를 유형화하여 접촉과 비접촉을 구별하면, 전화벨이 울리는 중에 휴대 전화에 손이 접촉할 경우 진동으로 전환된다.
- ⑤ 휴대 전화의 카메라로 촬영한 얼굴 영상들에서 색상값과 얼굴 형태 정보를 이용하여 얼굴과 얼굴이 아닌 것으로 분류하면, 사람이 움직여도 얼굴을 중심으로 촬영한다.

38. ㉡을 이용한 필기체 숫자 분류기의 구성도로 옳은 것은?



39. 문맥상 ㉠와 바꾸어 쓸 수 있는 한자어로 가장 적절한 것은?

- ① 취합(聚合)하여
- ② 융합(融合)하여
- ③ 조합(組合)하여
- ④ 규합(糾合)하여
- ⑤ 결합(結合)하여

[16~19] 다음 글을 읽고 물음에 답하시오.

인간의 신경 조직을 수학적으로 모델링하여 컴퓨터가 인간 처럼 기억·학습·판단할 수 있도록 구현한 것이 인공 신경망 기술이다. 신경 조직의 기본 단위는 뉴런인데, ㉠ 인공 신경망에서는 뉴런의 기능을 수학적으로 모델링한 퍼셉트론을 기본 단위로 사용한다.

㉢ 퍼셉트론은 입력값들을 받아들이는 여러 개의 ㉡ 입력 단자와 이 값을 처리하는 부분, 처리된 값을 내보내는 한 개의 출력 단자로 구성되어 있다. 퍼셉트론은 각각의 입력 단자에 할당된 ㉣ 가중치를 입력값에 곱한 값들을 모두 합하여 가중합을 구한 후, 고정된 ㉤ 임계치보다 가중합이 작으면 0, 그렇지 않으면 1과 같은 방식으로 ㉦ 출력값을 내보낸다.

이러한 퍼셉트론은 출력값에 따라 두 가지로만 구분하여 입력값들을 판정할 수 있을 뿐이다. 이에 비해 복잡한 판정을 할 수 있는 인공 신경망은 다수의 퍼셉트론을 여러 계층으로 배열하여 한 계층에서 출력된 신호가 다음 계층에 있는 모든 퍼셉트론의 입력 단자에 입력값으로 입력되는 구조로 이루어진다. 이러한 인공 신경망에서 가장 처음에 입력값을 받아들이는 퍼셉트론들을 입력층, 가장 마지막에 있는 퍼셉트론들을 출력층이라고 한다.

㉧ 어떤 사진 속 물체의 색깔과 형태로부터 그 물체가 사과인지 아닌지를 구별할 수 있도록 인공 신경망을 학습시키는 경우를 생각해 보자. 먼저 학습을 위한 입력값들 즉 학습 데이터를 만들어야 한다. 학습 데이터를 만들기 위해서는 사과 사진을 준비하고 사진에 나타난 특징인 색깔과 형태를 수치화해야 한다. 이 경우 색깔과 형태라는 두 범주를 수치화하여 하나의 학습 데이터로 묶은 다음, ‘정답’에 해당하는 값과 함께 학습 데이터를 인공 신경망에 제공한다. 이때 같은 범주에 속하는 입력값은 동일한 입력 단자를 통해 들어가도록 해야 한다. 그리고 사과 사진에 대한 학습 데이터를 만들 때에 정답인 ‘사과이다’에 해당하는 값을 ‘1’로 설정하였다면 출력값 ‘0’은 ‘사과가 아니다’를 의미하게 된다.

인공 신경망의 작동은 크게 학습 단계와 판정 단계로 나뉜다. 학습 단계는 학습 데이터를 입력층의 입력 단자에 넣어 주고 출력층의 출력값을 구한 후, 이 출력값과 정답에 해당하는 값의 차이가 줄어들도록 가중치를 갱신하는 과정이다. 어떤 학습 데이터가 주어지면 이때의 출력값을 구하고 학습 데이터와 함께 제공된 정답에 해당하는 값에서 출력값을 뺀 값 즉 오차 값을 구한다. 이 오차 값의 일부가 출력층의 출력 단자에서 입력층의 입력 단자 방향으로 되돌아가면서 각 계층의 퍼셉트론별로 출력 신호를 만드는 데 관여한 모든 가중치들에 더해지는 방식으로 가중치들이 갱신된다. 이러한 과정을 다양한 학습 데이터에 대하여 반복하면 출력값들이 각각의 정답 값에 수렴하게 되고 판정 성능이 좋아진다. 오차 값이 0에 근접하게 되거나 가중치의 갱신이 더 이상 이루어지지 않게 되면 학습 단계를 마치고 판정 단계로 전환한다. 이때 판정의 오류를 줄이기 위해서는 학습 단계에서 대상들의 변별적 특징이 잘 반영되어 있는 서로 다른 학습 데이터를 사용하는 것이 좋다.

16. 밑글에 따를 때, ㉠~㉦에 대한 설명으로 적절하지 않은 것은?

- ① ㉠은 ㉠의 기본 단위이다.
- ② ㉡는 ㉠을 구성하는 요소 중 하나이다.
- ③ ㉢가 변하면 ㉣도 따라서 변한다.
- ④ ㉤는 ㉦를 결정하는 기준이 된다.
- ⑤ ㉥가 학습하는 과정에서 ㉦는 ㉣의 변화에 영향을 미친다.

17. 밑글에 대한 이해로 적절하지 않은 것은?

- ① 퍼셉트론의 출력 단자는 하나이다.
- ② 출력층의 출력값이 정답에 해당하는 값과 같으면 오차 값은 0이다.
- ③ 입력층 퍼셉트론에서 출력된 신호는 다음 계층 퍼셉트론의 입력값이 된다.
- ④ 퍼셉트론은 인간의 신경 조직의 기본 단위의 기능을 수학적으로 모델링한 것이다.
- ⑤ 가중치의 갱신은 입력층의 입력 단자에서 출력층의 출력 단자 방향으로 진행된다.

18. 밑글을 바탕으로 ㉠에 대해 추론한 것으로 적절하지 않은 것은?

- ① 학습 데이터를 만들 때는 색깔이나 형태가 다른 사과와 사진을 선택하는 것이 좋겠군.
- ② 학습 데이터에 두 가지 범주가 제시되었으므로 입력층의 퍼셉트론은 두 개의 입력 단자를 사용하겠군.
- ③ 색깔에 해당하는 범주와 형태에 해당하는 범주를 분리하여 각각 서로 다른 학습 데이터로 만들어야 하겠군.
- ④ 가중치가 더 이상 변하지 않는 단계에 이르면 '사과'인지 아닌지를 구별하는 학습 단계가 끝났다고 볼 수 있겠군.
- ⑤ 학습 데이터를 만들 때 사과 사진의 정답에 해당하는 값을 0으로 설정하였다면, 출력층의 출력 단자에서 0 신호가 출력되면 '사과이다'로, 1 신호가 출력되면 '사과가 아니다'로 해석해야 되겠군.

19. 밑글을 바탕으로 <보기>를 이해한 내용으로 가장 적절한 것은? [3점]

—<보 기>—

아래의 [A]와 같은 하나의 퍼셉트론을 [B]를 이용해 학습 시키고자 한다.

[A]

- 입력 단자는 세 개(a, b, c)
- a, b, c의 현재의 가중치는 각각 $W_a=0.5$, $W_b=0.5$, $W_c=0.1$
- 가중합이 임계치 1보다 작으면 0을, 그렇지 않으면 1을 출력

[B]

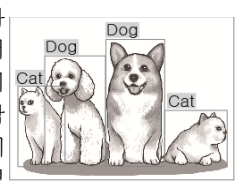
- a, b, c로 입력되는 학습 데이터는 각각 $I_a=1$, $I_b=0$, $I_c=1$
- 학습 데이터와 함께 제공되는 정답=1

- ① [B]로 학습시키기 위해서는 판정 단계를 먼저 거쳐야 하겠군.
- ② 이 퍼셉트론이 1을 출력한다면, 가중합이 1보다 작았기 때문이겠군.
- ③ [B]로 한 번 학습시키고 나면 가중치 W_a , W_b , W_c 가 모두 늘어나 있겠군.
- ④ [B]로 여러 차례 반복해서 학습시키면 퍼셉트론의 출력값은 0에 수렴하겠군.
- ⑤ [B]의 학습 데이터를 한 번 입력했을 때 그에 대한 퍼셉트론의 출력값은 1이겠군.

◆ 22년 7월 고3 10~13번

[10~13] 다음 글을 읽고 물음에 답하시오.

객체 탐지(Object Detection)란 사람, 동물, 사물 등 이미지에 있는 여러 대상의 위치를 찾아 각 대상의 크기에 맞는 경계 상자를 표시하고, 미리 학습된 객체 데이터를 바탕으로 그 경계 상자 안의 대상이 어떤 객체인지 판별하는 작업이다.



<객체 탐지의 예>

딥러닝 기반의 객체 탐지 모델은 ‘2단계 방식’과 ‘단일 단계 방식’으로 나눌 수 있다. 2단계 방식은 먼저 이미지에서 탐지할 객체가 있을 확률이 높은 곳을 추정한 후, 그 영역의 대상을 집중적으로 탐지하여 어떤 객체인지 판별하는 방식이다. 각 과정이 별도의 인공지능망을 통해 순차적으로 이루어지기 때문에 객체를 판별해 내는 정확도는 높지만, 처리하는 데이터가 많고 구조가 복잡하여 탐지 속도가 느리기 때문에 실시간으로 객체를 탐지하는 데는 어려움이 있다. 단일 단계 방식은 이 두 가지 과정이 하나의 인공지능망을 통해 동시에 이루어지는 방식인데, 가장 대표적인 알고리즘 모델로 YOLO(You Only Look Once)가 있다.

YOLO는 이미지가 입력되면 먼저 이미지를 $S \times S$ 개의 영역으로 나누고, 하나의 영역을 기준으로 경계 상자 N 개를 표시한다. 그리고 모든 영역마다 동일하게 N 개의 경계 상자를 표시하면서 각각의 경계 상자에 특정 객체가 존재할 확률도 예측한다. 이때 경계 상자의 개수가 많을수록 탐지 속도가 느려지기 때문에 설정할 수 있는 경계 상자의 수가 제한적인데 일반적으로 N 은 5 이하로 설정한다. 각 경계 상자의 데이터는 B_x , B_y , B_w , B_h , P_c 와 C 로 표시되는데 B_x , B_y 는 경계 상자의 중심점 좌표이며 B_w , B_h 는 폭과 높이이다. 그리고 P_c 는 해당 경계 상자에 어떤 객체가 존재할 확률값이고, C 는 그 객체가 특정 객체일 확률값이다. 이때 B_x , B_y 는 항상 기준이 되는 하나의 영역 안에 속해 있지만, 경계 상자의 크기는 영역의 크기와 상관없이 다양하게 표시된다. C 는 미리 학습된 m 가지 종류의 객체 데이터와 비교하여 각 객체일 확률을 표시한 값으로, 미리 학습된 객체의 가짓수에 따라 판별할 수 있는 객체의 가짓수가 결정되며, 그에 따라 C 의 개수도 결정된다. 하나의 이미지가 입력되면 이러한 방식으로 모든 영역별로 이미지에 있는 대상들을 확인하고 그 대상이 특정 객체일 확률값을 계산해서 총 ‘ $S \times S \times N(5+m)$ ’개의 데이터를 출력하게 된다.

이후 경계 상자에 객체가 존재할 확률값과 그것이 특정 객체일 확률값을 곱하여 해당 경계 상자에 특정 객체가 존재할 확률값인 ‘신뢰도 점수’를 구한다. 신뢰도 점수는 경계 상자의 위치와 객체의 판별이 얼마나 정확한지를 나타낸다. 모든 경계 상자들은 미리 학습된 객체의 가짓수만큼 신뢰도 점수를 가지

며, 이 중 가장 큰 값을 가지는 객체가 해당 경계 상자에서 탐지된 객체가 된다.

그런데 서로 다른 경계 상자에서 같은 종류의 객체가 탐지될 수 있다. 이때는 각 경계 상자가 하나의 대상에 중복되어 표시된 것인지, 서로 다른 대상에 표시된 것인지를 판단하여 이미지 속의 각 대상별로 가장 정확한 경계 상자 하나만 표시하는 과정을 거치는데, 이를 '비최대값 억제(NMS, Non-Max Suppression)'라고 한다. NMS는 두 경계 상자의 교집합을 합집합으로 나눈 값인 IoU를 기준으로 이루어진다. IoU 값은 두 경계 상자의 위치가 일치할수록 1에 가까운 값이 나오며, 이 값이 설정된 임계값보다 크면 두 경계 상자가 동일한 대상에 표시된 것으로 판단하고 둘 중 신뢰도 점수가 낮은 상자를 삭제한다. 그리고 IoU 값이 설정된 임계값보다 작으면 경계 상자가 서로 다른 대상에 표시된 것으로 판단하여 두 경계 상자 모두 그대로 둔다. 이러한 방법으로 한 가지 종류의 객체에 대해 그려진 모든 경계 상자들 중 가장 높은 신뢰도 점수를 가진 경계 상자를 기준으로 다른 경계 상자들을 하나씩 삭제해 나간다. 이후 IoU 값이 설정된 임계값보다 작아서 지워지지 않고 남겨진 경계 상자 중에서 가장 높은 신뢰도 점수를 가진 경계 상자를 다음 기준으로 정하여 동일한 과정을 반복한다. 그리고 이러한 과정을 다른 모든 대상에 표시된 경계 상자들에 대해서도 순차적으로 반복한다. 이렇게 해서 결국 이미지 속의 각 대상별로 가장 높은 신뢰도 점수를 가진 경계 상자 하나씩만 남게 된다.

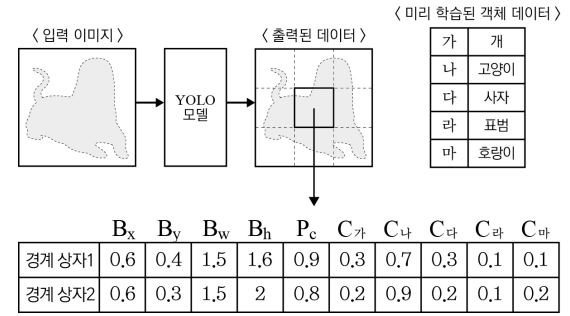
이런 원리로 인해 YOLO는 2단계 방식에 비해 탐지 속도가 매우 빨라서 자율 주행 자동차, 지능형 CCTV 등에 널리 사용되고 있다. 하지만 속도가 빠른 대신 ㉠새 떼와 같이 여러 물체가 한 영역 안에 모여 있는 경우 일부 대상을 탐지하지 못하는 한계점도 가지고 있다.

10. 윗글의 내용과 일치하지 않는 것은?

- ① 객체 탐지는 이미지에 있는 대상의 위치를 찾고 그 대상이 어떤 객체인지 판별하는 작업이다.
- ② 2단계 방식은 객체를 탐지하는 속도가 느려서 실시간 탐지에는 사용하기가 어렵다.
- ③ 이미지에 표시되는 경계 상자는 기준이 되는 영역의 크기에 따라 그 크기가 결정된다.
- ④ 신뢰도 점수는 경계 상자에 특정 객체가 존재할 확률값을 말하며 모든 경계 상자마다 존재한다.
- ⑤ 경계 상자가 표시되는 과정에서 하나의 대상에 여러 개의 경계 상자가 그려질 수도 있다.

11. 윗글을 참고할 때, <보기>에 대한 설명으로 적절하지 않은 것은?

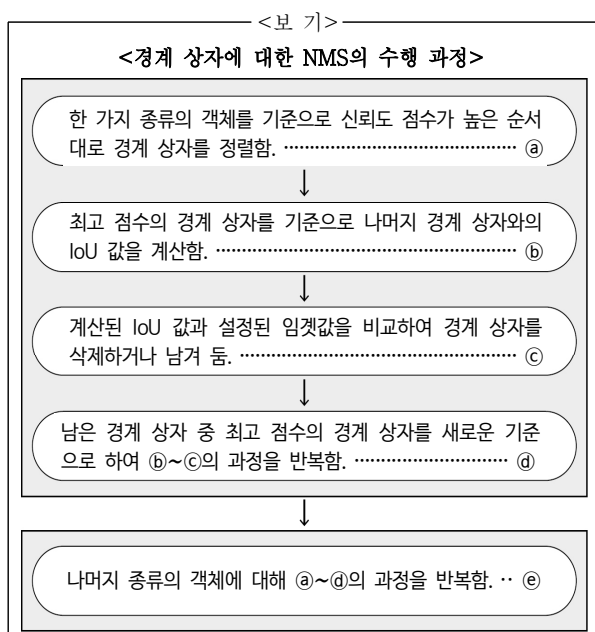
다음은 경계 상자의 수를 2로 설정한 YOLO 모델에 특정 이미지를 입력했을 때, 데이터가 출력되는 과정을 도식화하여 나타낸 것이다. 단, 입력된 이미지는 단일 객체에 대한 이미지이다.



- ① 입력된 이미지에서 탐지된 객체는 '고양이'일 가능성이 가장 높다.
- ② 경계 상자 1이 경계 상자 2보다 더 정확하게 객체를 탐지하였다.
- ③ 입력된 이미지의 전체 영역에 표시되는 경계 상자는 모두 18개이다.
- ④ 입력된 이미지에서 탐지할 수 있는 객체의 종류는 모두 다섯가지이다.
- ⑤ YOLO 모델이 이미지를 분석하여 출력하는 데이터는 모두 180개이다.

12. [A]를 바탕으로 <보기>를 이해한 내용으로 가장 적절한 것은?

[3점]



-
- ① ㉠의 대상이 되는 경계 상자의 신뢰도 점수는 이미지에 상관 없이 항상 일정하겠군.
 - ② ㉢에서 계산된 IoU 값이 0에 가까울수록 두 경계 상자는 중복되는 부분이 많겠군.
 - ③ ㉢의 과정에서 경계 상자가 삭제되지 않았다면 두 경계 상자가 동일한 대상에 표시된 경계 상자라고 판단한 것이겠군.
 - ④ ㉠의 과정은 하나의 특정 대상에 중복되어 표시된 여러 개의 경계 상자가 하나만 남을 때까지 반복되겠군.
 - ⑤ ㉢에서 새로운 기준이 되는 경계 상자는 이전 객체의 기준이 되었던 경계 상자와 동일한 대상에 그려져 있겠군.

13. ㉠의 이유로 가장 적절한 것은?

- ① 대상의 크기에 따라 해당 경계 상자에 존재할 확률값이 달라지기 때문에
- ② 객체에 대한 신뢰도 점수가 임계값보다 작아 경계 상자가 제거되기 때문에
- ③ 객체를 탐지할 때 미리 학습된 객체 데이터에 따라 객체를 판별하기 때문에
- ④ 객체를 탐지할 때 영역별로 탐지할 수 있는 객체의 수가 제한적이기 때문에
- ⑤ 객체를 탐지할 때 처리하는 데이터가 많고 알고리즘의 구조가 복잡하기 때문에

[10~13] 다음 글을 읽고 물음에 답하시오.

문장이나 영상, 음성을 만들어 내는 인공 지능 생성 모델 중 확산 모델은 영상의 복원, 생성 및 변환에 뛰어난 성능을 보인다. 확산 모델의 기본 발상은, 원본 이미지에 노이즈를 점진적으로 추가하였다가 그 노이즈를 다시 제거해 나가면 원본 이미지를 복원할 수 있다는 것이다. 노이즈는 불필요하거나 원하지 않는 값을 의미한다. 원하는 값만 들어 있는 원본 이미지에 노이즈를 단계별로 더하면 노이즈가 포함된 확산 이미지가 되고, 여러 단계를 거치면 결국 원본 이미지가 어떤 이미지였는지 전혀 알 수 없는 노이즈 이미지가 된다. 역으로, 단계별로 더해진 노이즈를 알 수 있다면 노이즈 이미지에서 원본 이미지를 복원할 수 있다. 확산 모델은 노이즈 생성기, 이미지 연산기, 노이즈 예측기로 구성되며, 순확산 과정과 역확산 과정 순으로 작동한다.

순확산 과정은 이미지에 노이즈를 추가하면서 노이즈 예측기를 학습시키는 과정이다. 첫 단계에서는, 노이즈 생성기에서 노이즈를 만든 후 이미지 연산기가 이 노이즈를 원본 이미지에 더해 노이즈가 포함된 확산 이미지를 출력한다. 다음 단계부터는 노이즈 생성기에서 만든 노이즈를 이전 단계에서 출력된 확산 이미지에 더한다. 이러한 단계를 충분히 반복하면 최종적으로 노이즈 이미지가 출력된다. 이때 더해지는 노이즈는 크기나 분포 양상 등 그 특성이 단계별로 다르다. 따라서 노이즈 예측기는 단계별로 확산 이미지를 입력받아 이미지에 포함된 노이즈의 특성을 추출하여 수치들로 표현하고, 이 수치들을 바탕으로 노이즈를 예측한다. 노이즈 예측기 내부의 이러한 수치들을 **잠재 표현**이라고 한다. 노이즈 예측기는 잠재 표현을 구하고 노이즈를 예측하는 방식을 학습한다.

노이즈 예측기의 학습 방법은 기계 학습 중에서 지도 학습에 해당한다. 지도 학습은 학습 데이터에 정답이 주어져 출력과 정답의 차이가 작아지도록 모델을 학습시키는 방법이다. 노이즈 예측기를 학습시킬 때는 노이즈 생성기에서 만들어 넣어 준 노이즈가 정답에 해당하며 이 노이즈와 예측된 노이즈 사이의 차이가 작아지도록 학습시킨다.

역확산 과정은 노이즈 이미지에서 노이즈를 제거하여 원본 이미지를 복원하는 과정이다. 노이즈를 제거하려면 이미지에 단계별로 어떤 특성의 노이즈가 더해졌는지 알아야 하는데 노이즈 예측기가 이 역할을 한다. 노이즈 이미지 또는 중간 단계에서의 확산 이미지를 노이즈 예측기에 입력하면 이미지에 포함된 노이즈의 특성을 추출하여 잠재 표현을 구하고 이를 바탕으로 노이즈를 예측한다. 이미지 연산기는 입력된 확산 이미지로부터 이 노이즈를 빼서 현 단계의 노이즈를 제거한 확산 이미지를 출력한다. 확산 이미지에 이런 단계를 반복하면 결국 노이즈가 대부분 제거되어 원본 이미지에 가까운 이미지만 남게 된다.

한편, 많은 종류의 이미지를 학습시킨 후 학습된 이미지의 잠재 표현에 고유 번호를 붙이면 역확산 과정에서 이미지를 선택하여 생성할 수 있다. 또한 잠재 표현의 수치들을 조정하면 다른 특성의 노이즈가 생성되어 여러 이미지를 혼합하거나 실재하지 않는 이미지를 만들어 낼 수도 있다.

10. 학생이 윗글을 읽은 방법으로 적절하지 않은 것은?

- ① 확산 모델이 지도 학습을 사용한다는 점에 주목하고, 지도 학습 방법이 확산 모델에 어떻게 적용되는지 확인하며 읽었다.
- ② 확산 모델이 두 가지 과정으로 이루어진다는 점에 주목하고, 두 과정 중 어느 과정이 선행되어야 하는지 살피며 읽었다.
- ③ 확산 모델에서 노이즈의 중요성을 파악하고, 사용되는 노이즈의 종류가 모델의 성능에 미치는 영향을 이해하며 읽었다.
- ④ 잠재 표현의 개념을 파악하고, 그 개념을 바탕으로 확산 모델이 노이즈를 예측하고 제거하는 원리를 이해하며 읽었다.
- ⑤ 확산 모델의 구성 요소를 파악하고, 그 구성 요소가 노이즈 처리 과정에서 어떤 기능을 하는지 확인하며 읽었다.

11. 윗글을 이해한 내용으로 가장 적절한 것은?

- ① 노이즈 생성기는 순확산 과정에서만 작동한다.
- ② 확산 모델에서의 학습은 역확산 과정에서 이루어진다.
- ③ 이미지 연산기와 노이즈 예측기는 모두 확산 이미지를 출력한다.
- ④ 노이즈 예측기를 학습시킬 때는 예측된 노이즈가 정답으로 사용된다.
- ⑤ 역확산 과정에서 단계가 반복될수록 출력되는 확산 이미지는 원본 이미지와의 유사성이 줄어든다.

12. **잠재 표현**에 대한 설명으로 적절하지 않은 것은?

- ① 잠재 표현의 수치들을 조정하면 여러 이미지를 혼합할 수 있다.
- ② 역확산 과정에서 잠재 표현이 다르면 예측되는 노이즈가 다르다.
- ③ 확산 모델의 학습에는 잠재 표현을 구하는 방식이 포함되어 있다.
- ④ 잠재 표현은 이미지에 더해진 노이즈의 크기나 분포 양상에 따라 다른 값들이 얻어진다.
- ⑤ 잠재 표현은 노이즈 예측기가 원본 이미지를 입력받아 노이즈의 특성을 추출한 결과이다.

13. 윗글을 바탕으로 <보기>를 이해한 내용으로 적절하지 않은 것은? [3점]

<보 기>

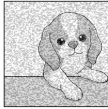
A 단계는 확산 모델 과정 중 한 단계이다. ㉠은 원본 이미지이고, ㉡은 확산 이미지 중의 하나이며, ㉢은 노이즈 이미지이다. (가)는 이미지가 A 단계로 입력되는 부분이고, (나)는 이미지가 A 단계에서 출력되는 부분이다.

(가) ⇨ A 단계 ⇨ (나)

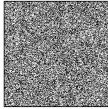
㉠



㉡



㉢



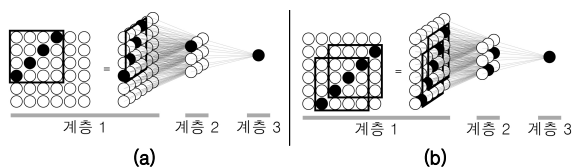
- ① (가)에 ㉠이 입력된다면, A 단계의 이미지 연산기에서는 ㉠에 노이즈를 더하겠군.
- ② (나)에 ㉢이 출력된다면, A 단계의 노이즈 생성기에서 생성된 노이즈가 이미지 연산기에서 확산 이미지에 더해졌겠군.
- ③ 순확산 과정에서 (가)에 ㉡이 입력된다면, A 단계의 노이즈 예측기에서 예측한 노이즈가 이미지 연산기에 입력되었겠군.
- ④ 역확산 과정에서 (가)에 ㉢이 입력된다면, A 단계의 이미지 연산기에서는 ㉢에서 노이즈를 빼겠군.
- ⑤ 역확산 과정에서 (나)에 ㉡이 출력된다면, A 단계의 노이즈 예측기에서 예측한 노이즈가 이미지 연산기에 입력되었겠군.

◆ 18년 3월 고2 24~28번

[24~28] 다음 글을 읽고 물음에 답하시오.

빛은 망막의 광수용기 세포에서 수용되어 전기 신호로 변환된 뒤, 뇌의 시각 피질로 전달된다. ㉠ 후벨과 위젤은 망막에 비춰진 빛에 대해 고양이 시각 피질 세포가 어떻게 반응하는지 실험하였다. 그들은 이를 통해 시각 피질 세포가 망막의 일정 영역 내 광수용기 세포들과 연결되어 있다는 사실을 알아냈다. 하나의 시각 피질 세포와 연결된 망막상의 일정 영역을 해당 시각 피질 세포의 '수용장'이라고 한다.

또한 이 실험을 통해 시각 피질이 하위의 '단순 세포'와 상위의 '복잡 세포'의 다층 구조로 구성되어 있다는 점이 밝혀졌다. 단순 세포와 복잡 세포 모두 각각의 수용장에 비친 특정한 각도를 가진 선분 모양의 빛에 활성화된다. 하지만 단순 세포가 수용장 내 특정 위치의 빛에만 활성화되는데 반해, 복잡 세포는 수용장이 단순 세포보다 넓고, 수용장에 비춰진 빛의 위치 변화에 관계없이 활성화된다. 이는 복잡 세포가 다수의 단순 세포들로부터 전기 신호를 전달받아 활성화되기 때문이다.

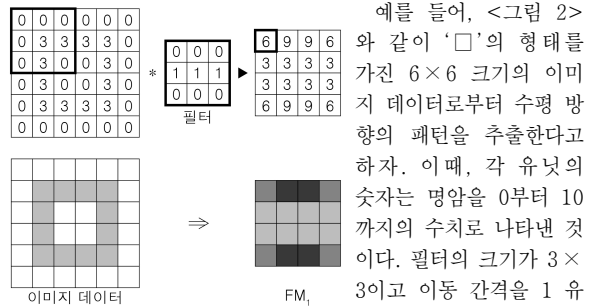


<그림 1>

<그림 1>은 이러한 시각 피질 세포들의 전기 신호 전달 과정을 다층 모형으로 나타낸 것이다. 모형의 각 층은 유닛들로 구성되는데, 계층 1의 각 유닛은 망막의 광수용기 세포에, 계층 2의 각 유닛은 단순 세포에, 계층 3의 유닛은 복잡 세포에 대응된다. 이때, 검은색 유닛은 해당 유닛이 활성화되었음을 의미하며, 계층 1의 사각형 영역은 계층 2의 활성화된 유닛의 수용장을 표시한 것이다. (a)와 (b)는 각각의 사선 패턴의 위치에 따른 각 유닛들의 활성화 상태를 나타낸 것이다. 계층 2의 각 유닛은 자신의 수용장 안의 특정한 위치에 특정한 각도의 사선 패턴이 입력되면 활성화된다. 계층 3의 유닛은 계층 2의 유닛 중에 하나라도 활성화되면 활성화된다.

'합성곱 신경망'은 이미지 인식(image recognition)*을 위해 만들어진 인공 신경망으로서, <그림 1>과 같은 다층 구조의 신경망 모형을 수학적으로 구조화한 것이다. 합성곱 신경망은

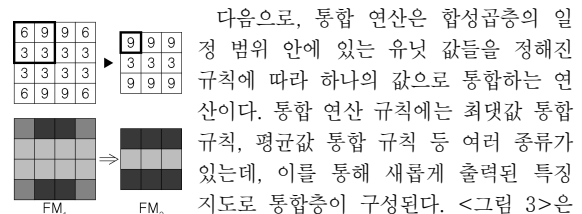
'합성곱층'과 '통합층'으로 구성되며, 이들은 각각 합성곱 연산과 통합 연산에 의해 출력된다. 먼저, 합성곱 연산은 특정한 크기의 [필터]가 이미지 데이터의 왼쪽 상단에서 오른쪽 하단까지 일정 간격으로 이동해 가며 이미지 데이터와 필터의 곱을 합산하는 과정이다. 이때 필터는 이미지 데이터의 국부 영역에 존재하는 특정한 기하학적 패턴을 검출하는 역할을 한다.



<그림 2>

필터가 왼쪽 상단에서 오른쪽 하단으로 한 칸씩 이동해 가면서 합성곱을 16번 연산하고 4×4 크기의 '특징 지도'(feature map, FM)가 출력된다. <그림 2>에서 특징 지도 FM₁의 가장 왼쪽 위 유닛 값 '6'은 이미지 데이터의 왼쪽 위 3×3의 영역과 필터와의 곱의 총합인 '0×0+0×0+0×0+0×1+3×1+3×1+0×0+3×0+0×0'의 연산을 통해 구해진 것이다.

이렇게 필터를 이용해 이미지 데이터에 합성곱 연산을 수행하면 필터의 특성에 맞게 강조된 특징 지도를 얻을 수 있다. <그림 2>는 합성곱 연산 결과 수평 방향의 패턴이 강조되고 데이터 크기는 6×6에서 4×4로 줄어 출력된 특징 지도를 보여 준다. 이때, 필터의 이동 간격이 크게 설정된다면 출력되는 특징 지도의 크기를 줄여 데이터 처리를 빠르게 할 수 있는 장점이 있지만, 이미지의 특징을 놓칠 가능성이 증가하게 되는 단점이 있다.



<그림 3>

<그림 2>의 FM₁을 2×2 범위로 최댓값 통합 규칙에 따라 통합 연산한 것이다. 이때, 통합 연산의 범위를 왼쪽 상단에서 오른쪽 하단까지 1 유닛 단위로 이동하도록 설정하면 3×3 크기의 새로운 특징 지도 FM₂가 출력된다.

합성곱 연산을 통해 이미지의 어떤 영역에 어떤 패턴이 있는지를 추출할 수 있으며, 다양한 필터를 통해 이를 반복하면 이미지 속 사물을 인식할 수 있다. 하지만 연산을 반복하는 과정에서 패턴의 위치 정보를 계속 유지하게 되는데, 이는 일반적으로 불필요한 정보이다. 왜냐하면, 합성곱 연산을 통해 출력된 특징 지도 내에서 서로 인접한 유닛들은 미세한 위치 정보만 다를 뿐, 거의 비슷한 패턴 정보를 담고 있기 때문이다. 이때, 통합 연산 수행은 합성곱 연산의 결과에서 위치 정보를 줄여 주는 역할을 한다.

합성곱 연산과 통합 연산을 통해 위치 정보는 축약되고 패턴 정보는 강조된 특징 지도가 출력된다. 그리고 이 특징 지도를 인공 지능 네트워크인 '전체 연결층'에 입력하여 이미지 인

식 결과를 출력할 수 있다. 또한 입력된 이미지가 많아질수록 인공 신경망의 기계 학습을 통해 합성곱 신경망이 스스로 필터의 수치를 갱신함으로써 이미지 인식의 정확성이 높아지게 된다. 하지만 합성곱 연산 및 통합 연산의 횟수, 필터의 크기 및 이동 간격, 통합 연산 규칙 등은 초기 설정 값이 계속 유지되므로 이를 고려하여 합성곱 신경망을 설계해야 한다. 최근 인공 지능 기술이 발전함에 따라 합성곱 신경망은 사진 자동 분류, 필기 인식 등 다양한 영역으로 확장되고 있다.

* 이미지 인식: 이미지 속 사물이 무엇인지를 알아내는 것.

24. 윗글에 대한 이해로 적절하지 않은 것은?

- ① 통합 연산은 합성곱층의 일정 범위 내의 값들을 하나의 값으로 통합하는 기능을 한다.
- ② 시각 피질의 복잡 세포는 단순 세포로부터 전달받은 전기 신호를 전체 연결층에 전달한다.
- ③ 시각 피질의 단순 세포는 수용장 내에 비취진 특정 각도의 선분 모양의 빛에 활성화된다.
- ④ 합성곱 신경망으로 이미지를 인식하려면 특정 지도에 특정 패턴에 대한 정보가 담겨 있어야 한다.
- ⑤ 합성곱 신경망은 합성곱 연산과 통합 연산을 통해 이미지의 패턴 정보가 강조된 특정 지도를 추출한다.

25. <보기>는 ㉠을 재구성한 실험에 대한 설명이다. <보기>와 윗글의 <그림 1>을 이해한 것으로 적절하지 않은 것은? [3점]

< 보 기 >

다양한 빛 자극에 대해 시각 피질 세포가 어떻게 반응하는지 알기 위해, 선분 모양의 빛을 고향이의 망막에 비춘다. 이때, 빛의 각도는 각도 ㉠과 ㉡로, 빛이 비추어지는 수용장 내 위치는 위치 ㉢과 ㉣로 각각 다르게 한다. 그 결과 세포 A와 B는 서로 다른 반응을 보였다.

(단, 세포 A와 B는 서로 다른 시각 피질 세포이며, 망막의 특정 영역을 수용장으로 공유한다.)

실험			실험 결과	
	빛의 각도	빛의 위치	세포 A	세포 B
자극 1	㉠	㉢	○	○
자극 2	㉠	㉣	○	×
자극 3	㉡	㉢	×	×
자극 4	㉡	㉣	×	×

(○: 활성화, ×: 비활성화)

- ① '자극 1'의 실험 결과를 고려하면, '세포 A'와 '세포 B'가 반응하는 빛의 각도는 같겠군.
- ② '자극 1'과 '자극 2'의 실험 결과를 고려하면, '세포 A'의 수용장이 '세포 B'의 수용장보다 더 넓겠군.
- ③ '자극 1'과 '자극 3'의 실험 결과를 비교하면, '세포 A'는 각도 ㉡의 빛에는 반응하지 않겠군.
- ④ '세포 A'는 <그림 1>의 '계층 3'의 유닛에, '세포 B'는 '계층 2'의 유닛에 해당하겠군.
- ⑤ '자극 1'과 '자극 2'의 실험 결과는 <그림 1>의 (a)에, '자극 3'과 '자극 4'의 실험 결과는 (b)에 해당하겠군.

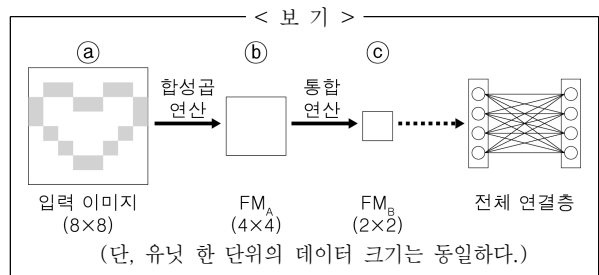
26. [필터]에 대한 이해로 적절하지 않은 것은?

- ① 합성곱 연산을 수행하면 필터의 특성이 반영된 특징 지도가 출력된다.
- ② 필터의 기능은 이미지 데이터에서 특정한 기하학적 패턴을 검출하는 것이다.
- ③ 적절한 필터를 통해 합성곱 연산을 반복하여 이미지 속 사물을 인식할 수 있다.
- ④ 필터의 크기와 이동 간격의 비율은 합성곱 신경망에 의해 자동적으로 변화된다.
- ⑤ 필터의 매개를 통해 이미지 속 사물의 패턴에 대한 정보가 합성곱층에 반영된다.

27. [가]를 고려할 때, '통합 연산'을 수행하는 이유로 적절한 것은?

- ① 통합 연산 수행 이전과 이후, 이미지 속 사물에 대한 인식 결과가 달라지기 때문이다.
- ② 통합층의 각 유닛에 담긴 정보는 합성곱층의 각 유닛에 담긴 정보와 관련이 없기 때문이다.
- ③ 이미지 속 사물의 위치 정보를 표시하기 위한 추가적인 합성곱 연산이 필요하기 때문이다.
- ④ 합성곱 연산을 수행한 결과에 이미지 인식에는 불필요한 위치 정보가 포함되어 있기 때문이다.
- ⑤ 통합 연산은 합성곱층에 포함된 이미지 속 사물의 패턴 정보를 추출하는 역할을 하기 때문이다.

28. <보기>는 '♡' 모양의 디지털 이미지를 인식하는 과정의 일부를 도식화한 것이다. 이에 대해 이해한 것으로 적절하지 않은 것은?



- ① ㉡의 데이터 크기는 ㉠에 비해 작겠군.
- ② 필터의 이동 간격을 1 유닛 단위로 설정했다면 ㉡를 출력하기 위해 5×5 필터가 사용되었겠군.
- ③ 2×2 범위로 평균값 통합을 통해 ㉡를 출력했다면, ㉡의 데이터 크기는 ㉡의 25%로 감소하였겠군.
- ④ 2×2 범위로 최댓값 통합 규칙을 사용하여 ㉡를 통합 연산한 경우, 해당 범위의 유닛 값들 중 최댓값이 ㉡의 하나의 유닛 값으로 도출되었겠군.
- ⑤ ㉡에서 ㉡를 출력하기 위한 통합 연산에는 '♡' 모양의 특징을 검출할 수 있는 필터가 적용되었겠군.