



(19) 대한민국특허청(KR)
(12) 공개특허공보(A)

(11) 공개번호 10-2024-0102684
(43) 공개일자 2024년07월03일

(51) 국제특허분류(Int. Cl.)
G06N 3/098 (2023.01) G06N 3/063 (2023.01)
(52) CPC특허분류
G06N 3/098 (2023.01)
G06N 3/063 (2023.01)
(21) 출원번호 10-2022-0184961
(22) 출원일자 2022년12월26일
심사청구일자 없음

(71) 출원인
서강대학교산학협력단
서울특별시 마포구 백범로 35 (신수동, 서강대학교)
(72) 발명자
박성용
서울특별시 은평구 수색로 216, DMC SK뷰 아파트
105동 1501호
맹산하
서울특별시 성북구 송인로 50 래미안길음센터퍼스
114-1401
(74) 대리인
이지연

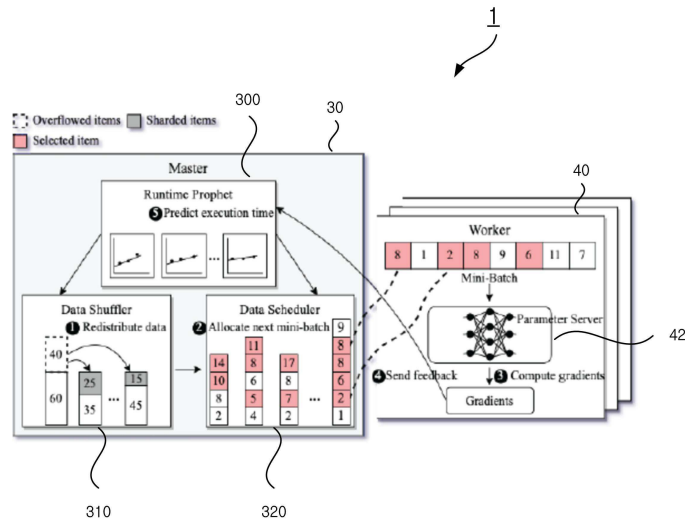
전체 청구항 수 : 총 11 항

(54) 발명의 명칭 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 시스템 및 그 방법

(57) 요약

본 발명은 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 시스템에 관한 것이다. 상기 분산 딥러닝 가속화 시스템은, 각 노드의 학습 성능을 평가하고, 각 노드의 학습 성능에 따라 각 노드에서의 데이터당 학습 시간을 예측하고, 예측된 각 노드의 데이터당 예상 학습 시간을 이용하여 노드들의 학습시간이 균일하도록 데이터를 스케줄링하고, 스케줄링된 데이터를 워커들에게 할당하는 마스터; 및 상기 마스터로부터 할당된 데이터들을 이용하여 노드에서 딥러닝 모델을 학습시키고, 학습이 완료될 때 마다 각 노드의 학습 데이터의 크기 및 학습 시간으로 이루어진 피드백을 마스터로 전송하는 하나 또는 둘 이상의 워커들;을 구비하여, 학습 데이터의 불균형을 최소화시켜 분산 딥러닝의 학습 속도를 가속화시킨다.

대표도 - 도3



(52) CPC특허분류
Z05S 12/00 (2022.08)

이 발명을 지원한 국가연구개발사업

과제고유번호	1711152712
과제번호	2017-0-01628-006
부처명	과학기술정보통신부
과제관리(전문)기관명	정보통신기획평가원
연구사업명	정보통신방송혁신인재양성
연구과제명	적응형 블록체인 플랫폼 기술 개발 및 전문 인력 양성
기 여 율	1/1
과제수행기관명	서강대학교산학협력단
연구기간	2022.01.01 ~ 2022.12.31

명세서

청구범위

청구항 1

컴퓨팅 시스템으로 이루어진 복수 개의 노드들로 구성된 클러스터를 이용하여 딥러닝 모델을 분산 학습시키는 분산 딥러닝 가속화 시스템에 있어서,

각 노드의 학습 성능을 평가하고, 각 노드의 학습 성능에 따라 각 노드에서의 데이터당 학습 시간을 예측하고, 예측된 각 노드의 데이터당 학습 시간을 이용하여 노드들의 학습시간이 동일하도록 데이터를 스케줄링하고, 스케줄링된 데이터를 워커들에게 할당하는 마스터; 및

상기 마스터로부터 할당된 데이터들을 이용하여 노드에서 딥러닝 모델을 학습시키고, 학습이 완료될 때 마다 각 노드의 학습 데이터의 크기 및 학습 시간으로 이루어진 피드백을 마스터로 전송하는 하나 또는 둘 이상의 워커들;

을 구비하여, 학습 데이터의 불균형을 최소화시켜 분산 딥러닝의 학습 속도를 가속화시키는 것을 특징으로 하는 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 시스템.

청구항 2

제1항에 있어서, 상기 마스터는,

워커로부터 전송된 학습한 데이터의 크기와 학습 시간의 집합들을 이용하여 노드의 성능 지표를 구하고, 성능 지표를 이용하여 각 노드의 성능을 평가하는 노드 성능 평가 모듈; 및

학습 데이터의 크기와 노드의 성능 지표를 이용하여, 각 노드의 데이터 당 예상 학습 시간을 계산하고, 각 노드의 데이터당 예상 학습 시간을 이용하여 각 노드가 학습할 데이터를 할당하는 데이터 스케줄링 모듈;

을 구비하여, 모든 노드의 예상 학습 시간이 균등하게 분포되도록, 각 노드에 할당되는 학습 데이터를 스케줄링하는 것을 특징으로 하는 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 시스템.

청구항 3

제1항에 있어서, 상기 마스터는

한 세대의 학습이 완료되면, 다음 세대에서 각 노드의 총 학습 시간이 균일하도록, 각 노드가 다음 세대에서 학습할 데이터를 섞어 재배치하는 데이터 셔플링 모듈;을 더 구비하는 것을 특징으로 하는 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 시스템.

청구항 4

제2항에 있어서, 상기 데이터 스케줄링 모듈은,

초기 배치 학습을 위하여 처음으로 할당하는 데이터는 무작위로 데이터를 선택하여 노드에게 할당하고,

그 후의 배치 학습에는, 노드 성능 평가 모듈에 의해 제공된 학습 데이터의 크기와 노드의 성능 지표를 이용하여, 각 노드의 데이터 당 예상 학습 시간을 계산하고, 각 노드의 데이터당 예상 학습 시간을 이용하여, 모든 노드의 예상 학습 시간이 균등하게 분포되도록, 각 노드에 할당되는 학습 데이터를 스케줄링하는 것을 특징으로 하는 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 시스템.

청구항 5

제1항에 있어서, 상기 마스터는

각 노드에서 학습한 데이터의 크기에 비례하도록 학습률을 조정하여 학습률 보정을 하는 학습률 보정 모듈;을 더 구비하는 것을 특징으로 하는 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 시스템.

청구항 6

컴퓨팅 시스템에서 실행되는 마스터와 워커로 이루어진 프로그램으로 구현된 딥러닝 모델을 분산 학습시키는 분산 딥러닝 가속화 방법에 있어서,

(a) 마스터에 의해, 각 노드의 학습 성능을 평가하고, 각 노드의 학습 성능에 따라 각 노드에서의 데이터당 학습 시간을 예측하고, 예측된 각 노드의 데이터당 학습 시간을 이용하여 노드들의 학습시간이 동일하도록 데이터를 스케줄링하고, 스케줄링된 데이터를 워커들에게 할당하는 단계; 및

(b) 워커에 의해, 상기 마스터로부터 할당된 데이터들을 이용하여 노드에서 딥러닝 모델을 학습시키고, 학습이 완료될 때 마다 각 노드의 학습 데이터의 크기 및 학습 시간으로 이루어진 피드백을 마스터로 전송하는 단계;

를 구비하여, 학습 데이터의 불균형을 최소화시켜 분산 딥러닝의 학습 속도를 가속화시키는 것을 특징으로 하는 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 방법.

청구항 7

제6항에 있어서, 상기 (a) 단계는,

(a1) 워커로부터 전송된 학습한 데이터의 크기와 학습 시간의 집합들을 이용하여 노드의 성능 지표를 구하고, 성능 지표를 이용하여 각 노드의 성능을 평가하는 단계; 및

(a2) 학습 데이터의 크기와 노드의 성능 지표를 이용하여, 각 노드의 데이터 당 예상 학습 시간을 계산하고, 각 노드의 데이터당 예상 학습 시간을 이용하여 각 노드가 학습할 데이터를 할당하는 단계;

를 구비하여, 모든 노드의 예상 학습 시간이 균등하게 분포되도록, 각 노드에 할당되는 학습 데이터를 스케줄링하는 것을 특징으로 하는 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 방법.

청구항 8

제7항에 있어서, 상기 (a) 단계는

(a3) 한 세대의 학습이 완료되면, 다음 세대에서 각 노드의 총 학습 시간이 균일하도록, 각 노드가 다음 세대에서도 학습할 데이터를 섞어 재배치하는 데이터 셔플링 단계;를 더 구비하는 것을 특징으로 하는 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 방법.

청구항 9

제7항에 있어서, 상기 (a2) 단계는,

초기 배치 학습을 위하여 처음으로 할당하는 데이터는 무작위로 데이터를 선택하여 노드에게 할당하고,

그 후의 배치 학습에는, 노드 성능 평가 모듈에 의해 제공된 학습 데이터의 크기와 노드의 성능 지표를 이용하여, 각 노드의 데이터 당 예상 학습 시간을 계산하고, 각 노드의 데이터당 예상 학습 시간을 이용하여, 모든 노드의 예상 학습 시간이 균등하게 분포되도록, 각 노드에 할당되는 학습 데이터를 스케줄링하는 것을 특징으로 하는 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 방법.

청구항 10

제7항에 있어서, 상기 (a) 단계는,

(a4) 각 노드에서 학습한 데이터의 크기에 비례하도록 학습률을 조정하여 학습률 보정을 하는 단계;를 더 구비하는 것을 특징으로 하는 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 방법.

청구항 11

제6항에 있어서, 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 방법은,

(c) 마스터가 학습 데이터에 대한 불균형 정도(Data Imbalance Factor)를 아래의 수학적식에 따라 측정하는 단계;를 더 구비하는 것을 특징으로 하는 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 방법.

발명의 설명

기술 분야

- [0001] 본 발명은 분산 딥러닝 가속화 기법에 관한 것으로서, 더욱 구체적으로는 학습 데이터 크기의 불균형을 고려하여, 클러스터를 구성하는 노드들이 동일한 학습 시간을 갖도록 데이터 스케줄링 및 서플링하는 정책을 제공함으로써, straggler 문제를 해결하고 딥러닝 모델의 수렴을 가속화시킬 수 있도록 하는 분산 딥러닝 가속화 시스템 및 그 방법에 관한 것이다.

배경 기술

- [0002] 딥러닝(Deep Learning)은 컴퓨터 비전, 자연어 처리 등 다양한 분야에서 좋은 성능을 보이고 있다. 딥러닝은 심층 신경망 모델을 구축한 후 학습 데이터셋을 입력으로 하여 학습 알고리즘으로 딥러닝 모델을 학습시킴으로써 수행된다. 최근, 더욱 우수한 성능의 딥러닝 모델을 얻기 위하여 학습 데이터 셋과 모델의 크기가 점차 커짐에 따라, 충분히 빠른 시간 내에 학습을 완료하기가 어려워졌다.
- [0003] 이를 위해, 학습 데이터셋을 여러 노드에 분산시켜 병렬적으로 학습을 수행하는 데이터 병렬화 방식의 분산 딥러닝이 제안되었다. 또한, 딥러닝은 계산 집약적이고 오랜 시간이 걸리는 작업이기 때문에, 학습의 가속화를 위해 GPU 클러스터를 사용한 분산 딥러닝이 널리 사용되고 있다.
- [0004] 분산 딥러닝 중에서도, 딥러닝 모델의 정확도를 중요시하는 응용들은 매개 변수 서버를 두고 딥러닝 모델의 정보를 동기화하는 방식으로 딥러닝 모델을 학습시킨다.
- [0005] 도 1은 종래의 기술에 따라, 4 개의 노드로 구성된 클러스터에서 데이터 병렬화 방식을 사용하여 학습을 병렬화하는 동기식 분산 딥러닝의 구조를 도시한 모식도이다. 도 1에 따른 방식은 입력 데이터에 대한 딥러닝 모델의 오차를 계산하고 역전파 과정을 통해 계산된 기울기를 각 노드의 매개 변수 서버들이 공유하여 딥러닝 모델의 매개 변수 정보를 동기화하는 과정을 반복하는 방식이다.
- [0006] 이러한 동기식 분산 딥러닝은 배치마다 동기화를 위해 모든 노드의 학습이 끝나기를 기다려야 하는데, 이 중 학습이 가장 오래 걸리는 노드를 straggler라 하며, straggler에 의해 학습이 지연되는 현상을 straggler 문제라 한다. 이와 같이, 분산 딥러닝에서의 학습의 지연 현상은 학습이 가장 느린 straggler에 의해 큰 영향을 받게 된다. 하지만, 기존의 딥러닝 시스템들은 이를 고려하지 않고 설계되었기 때문에, 노드 간 학습 시간의 차이로 인해 학습 속도와 확장성이 저하될 뿐만 아니라 성능 향상에도 한계가 발생하였다.
- [0007] 이와 같이, 분산 딥러닝은 straggler 문제에 의해 속도 개선에 한계가 존재하여, 이러한 straggler 문제를 해결하기 위한 다양한 연구들이 제안되고 있다. 제안된 해결 방안 중 하나는, N개의 노드에 추가적인 b개의 백업 노드를 사용하여 학습이 가장 빨리 완료된 N개 노드의 결과만 매개 변수 서버에 반영하는 방식이다. 다른 해결 방안은, 모든 노드에서 고정된 시간만큼만 학습을 수행하는 방식이다. 또 다른 해결 방안은, 클러스터의 이중 노드 간 성능 차이를 고려하여 각 노드의 성능에 비례하도록 배치 데이터의 크기를 동적으로 조절하는 방식이다.
- [0008] 위에서 언급한 선행 연구들은 이미지와 같이 고정된 크기의 학습 데이터만으로 학습을 수행하는 상황을 가정하였다. 하지만, 오늘날의 딥러닝은 학습 데이터셋이 크고 다양해짐에 따라, 영상, 음성, 문장, 신호 등과 같이 전처리 후에도 각 학습 데이터의 크기가 상이한 불균형 데이터셋으로 딥러닝을 많이 수행하고 있다. 이러한 불균형 데이터셋은 전처리후에도 각 학습 데이터의 크기가 달라 학습에 소요되는 시간이 서로 다르기 때문에 straggler 문제를 심화시켜 학습을 더욱 지연시키게 되는 문제점이 발생한다.

선행기술문헌

특허문헌

- [0009] (특허문헌 0001) 한국공개특허공보 제 10-2021-0047626호
(특허문헌 0002) 한국공개특허공보 제 10-2019-0021877호
(특허문헌 0003) 한국공개특허공보 제 10-2022-0031629호

(특허문헌 0004) 한국공개특허공보 제 10-2011-0165584호

발명의 내용

해결하려는 과제

- [0010] 전술한 문제점을 해결하기 위하여, 본 발명은 학습 데이터 크기의 불균형을 고려하여, 클러스터를 구성하는 노드들이 동일한 학습 시간을 갖도록 데이터 스케줄링 및 서플링하는 정책을 제공함으로써, 딥러닝 모델의 학습을 가속화시킬 수 있도록 하는 분산 딥러닝 가속화 시스템 및 그 방법을 제공하는 것을 목적으로 한다.

과제의 해결 수단

- [0011] 전술한 기술적 과제를 달성하기 위한 본 발명의 제1 특징에 따른 분산 딥러닝 가속화 시스템은, 컴퓨팅 시스템으로 이루어진 복수 개의 노드들로 구성된 클러스터를 이용하여 딥러닝 모델을 분산 학습시키는 분산 딥러닝 가속화 시스템에 관한 것으로서, 각 노드의 학습 성능을 평가하고, 각 노드의 학습 성능에 따라 각 노드에서의 데이터당 학습 시간을 예측하고, 예측된 각 노드의 데이터당 학습 시간을 이용하여 노드들의 학습 시간이 균일하도록 데이터를 스케줄링하고, 스케줄링된 데이터를 워커에게 할당하는 마스터; 및 상기 마스터로부터 할당된 데이터들을 이용하여 노드에서 딥러닝 모델을 학습시키고, 학습이 완료될 때 마다 각 노드의 학습 데이터의 크기 및 학습 시간으로 이루어진 피드백을 마스터로 전송하는 하나 또는 둘 이상의 워커들;을 구비하여, 학습 데이터의 불균형을 최소화시켜 딥러닝 모델의 학습을 가속화시킨다.
- [0012] 전술한 본 발명의 제1 특징에 따른 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 시스템에 있어서, 상기 마스터는, 워커로부터 전송된 학습한 데이터의 크기와 학습 시간의 집합들을 이용하여 노드의 성능 지표를 구하고, 성능 지표를 이용하여 각 노드의 성능을 평가하는 노드 성능 평가 모듈; 및 학습 데이터의 크기와 노드의 성능 지표를 이용하여, 각 노드의 데이터 당 예상 학습 시간을 계산하고, 각 노드의 데이터당 예상 학습 시간을 이용하여 각 노드가 학습할 데이터를 할당하는 데이터 스케줄링 모듈;을 구비하여, 모든 노드의 예상 학습 시간이 균등하게 분포되도록, 각 노드에 할당되는 학습 데이터를 스케줄링하는 것이 바람직하다.
- [0013] 전술한 본 발명의 제1 특징에 따른 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 시스템에 있어서, 상기 마스터는, 한 세대의 학습이 완료되면, 다음 세대에서 각 노드의 총 학습 시간이 균일하도록, 각 노드가 다음 세대에서 학습할 데이터를 섞어 재배치하는 데이터 서플링 모듈;을 더 구비하는 것이 바람직하다.
- [0014] 전술한 본 발명의 제1 특징에 따른 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 시스템에 있어서, 상기 데이터 스케줄링 모듈은, 학습의 무작위성을 위하여 처음으로 할당하는 데이터는 무작위로 데이터를 선택하여 노드에게 할당하고, 그 후의 배치 데이터에는, 노드 성능 평가 모듈에 의해 제공된 학습 데이터의 크기와 노드의 성능 지표를 이용하여, 각 노드의 데이터 당 예상 학습 시간을 계산하고, 각 노드의 데이터당 예상 학습 시간을 이용하여, 모든 노드의 예상 학습 시간이 균등하게 분포되도록, 각 노드에 할당되는 학습 데이터를 스케줄링하는 것이 바람직하다.
- [0015] 전술한 본 발명의 제1 특징에 따른 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 시스템에 있어서, 상기 마스터는 각 노드에서 학습한 데이터의 크기에 비례하도록 학습률을 조정하여 학습률 보정을 하는 학습률 보정 모듈;을 더 구비하는 것이 바람직하다.
- [0016] 본 발명의 제2 특징에 따른 분산 딥러닝 가속화 방법은, 컴퓨팅 시스템에서 실행되는 마스터와 워커로 이루어진 프로그램으로 구현된 딥러닝 모델을 분산 학습시키는 분산 딥러닝 가속화 방법에 관한 것으로서, (a) 마스터에 의해, 각 노드의 학습 성능을 평가하고, 각 노드의 학습 성능에 따라 각 노드에서의 데이터당 학습 시간을 예측하고, 예측된 각 노드의 데이터당 학습 시간을 이용하여 노드들의 학습시간이 동일하도록 데이터를 스케줄링하고, 스케줄링된 데이터를 워커들에게 할당하는 단계; 및 (b) 워커에 의해, 상기 마스터로부터 할당된 데이터들을 이용하여 노드에서 딥러닝 모델을 학습시키고, 학습이 완료될 때 마다 각 노드의 학습 데이터의 크기 및 학습 시간으로 이루어진 피드백을 마스터로 전송하는 단계;를 구비하여, 학습 데이터의 불균형을 최소화시켜 분산 딥러닝의 학습 속도를 가속화시킨다.
- [0017] 전술한 본 발명의 제2 특징에 따른 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 방법에 있어서, 상기 (a) 단계는, (a1) 워커로부터 전송된 학습한 데이터의 크기와 학습 시간의 집합들을 이용하여 노드의 성능 지표를 구하고, 성능 지표를 이용하여 각 노드의 성능을 평가하는 단계; 및 (a2) 학습 데이터의 크기와 노드의 성능

지표를 이용하여, 각 노드의 데이터 당 예상 학습 시간을 계산하고, 각 노드의 데이터당 예상 학습 시간을 이용하여 각 노드가 학습할 데이터를 할당하는 단계;를 구비하여, 모든 노드의 예상 학습 시간이 균등하게 분포되도록, 각 노드에 할당되는 학습 데이터를 스케줄링하는 것이 바람직하다.

[0018] 전술한 본 발명의 제2 특징에 따른 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 방법에 있어서, 상기 (a) 단계는 (a3) 한 세대의 학습이 완료되면, 다음 세대에서 각 노드의 총 학습 시간이 균일하도록, 각 노드가 다음 세대에서 학습할 데이터를 섞어 재배치하는 데이터 셔플링 단계;를 더 구비하는 것이 바람직하다.

[0019] 전술한 본 발명의 제2 특징에 따른 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 방법에 있어서, 상기 (a2) 단계는, 학습의 무작위성을 위하여 처음으로 할당하는 데이터는 무작위로 데이터를 선택하여 노드에게 할당하고, 그 후의 배치 데이터에는, 노드 성능 평가 모듈에 의해 제공된 학습 데이터의 크기와 노드의 성능 지표를 이용하여, 각 노드의 데이터 당 예상 학습 시간을 계산하고, 각 노드의 데이터당 예상 학습 시간을 이용하여, 모든 노드의 예상 학습 시간이 균등하게 분포되도록, 각 노드에 할당되는 학습 데이터를 스케줄링하는 것이 바람직하다.

[0020] 전술한 본 발명의 제2 특징에 따른 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 방법에 있어서, 상기 (a) 단계는, (a4) 각 노드에서 학습한 데이터의 크기에 비례하도록 학습물을 조정하여 학습물 보정을 하는 단계;를 더 구비하는 것이 바람직하다.

발명의 효과

[0021] 본 발명에 따른 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 시스템 및 방법은, 각 노드의 데이터당 예상 학습 시간을 계산하고 이를 이용하여 노드의 학습 시간들이 균일하도록 데이터를 스케줄링함으로써, 데이터의 크기가 불균형한 데이터 셋을 사용하여 분산 학습을 하더라도, straggler 문제를 해결할 수 있게 되고, 그 결과 straggler에 의한 학습 지연이 최소화 되어 딥러닝을 더욱 가속화시킬 수 있게 된다.

도면의 간단한 설명

[0022] 도 1은 종래의 기술에 따라, 4대의 노드로 구성된 클러스터에서 데이터 병렬화 방식을 사용하여 학습을 병렬화하는 동기식 분산 딥러닝의 구조를 도시한 모식도이다.

도 2는 DIF에 따른 딥러닝 모델별 학습 시간 증가량을 도시한 그래프이다.

도 3은 본 발명의 바람직한 실시예에 따른 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 시스템을 도시한 구조도이다.

도 4는 본 발명의 바람직한 실시예에 따른 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 시스템에 있어서, 전체 동작을 순차적으로 도시한 흐름도이다.

도 5는 DIF와 워크로드에 따른 학습 시간을 비교한 그래프이다.

도 6은 워크로드에 따른 모델 수렴 속도를 비교한 그래프이다.

발명을 실시하기 위한 구체적인 내용

[0023] 본 발명은 학습 데이터 셋으로부터 데이터 불균형 정도를 측정하고, 이에 따른 딥러닝 모델의 학습 시간 증가량을 분석하여, 데이터 불균형 정도와 학습 시간의 상관 관계를 수학적으로 모델링하고, 이를 통해 각 노드에서 학습한 데이터의 크기와 학습 시간으로부터 각 노드의 성능을 평가하여 데이터당 학습 시간을 예측하고 학습의 부하를 조절하여 straggler 문제를 최소화하는 기법 및 학습한 데이터의 크기에 비례하여 학습물을 보정하는 방안을 제안하고자 한다.

[0024] 먼저, 본 발명에서 사용되는 DIF와 학습 시간의 상관 관계에 대하여 설명한다. 데이터 불균형 정도를 나타내는 DIF(Data Imbalance Factor)는 학습 데이터셋의 크기 집합 D 에 대한 표준 편차로 정의되며, 수학식 1로 표현된다.

수학식 1

$$DIF = \sqrt{\frac{1}{|D|-1} \sum_{d \in D} (d - \mu(D))^2}$$

여기서, D는 학습 데이터 셋의 크기 집합이며, |D|는 D의 데이터 샘플들의 개수이며, $\mu()$ 는 평균 함수이다.

실제 데이터셋에서 데이터의 크기가 얼마나 불균형한지 측정하기 위하여 범용 영상 데이터셋인 UCF101, Charades, DAVIS, HACS, COIN, Youtube Bounding Boxes, Sports-1M 으로 데이터 불균형 정도를 측정하였다. 이때, 각 데이터는 학습 시 전처리되어 첫번째 차원을 제외한 다른 모든 차원은 고정된 크기를 갖는 텐서로 변환되므로 영상의 길이만이 학습시 데이터의 불균형 정도를 결정한다. 표 1은 전술한 일곱 가지의 데이터 셋에서 측정한 DIF 를 나타낸다.

표 1

학습 데이터 셋	DIF
UCF101	3.73
Charades	9.45
DAVIS	18.36
HACS	64.13
COIN	70.02
YouTube Bounding Boxes	212.64
Sports-1M	539.75

불균형 데이터 셋으로 분산 딥러닝을 수행할 경우, 각 노드에 할당되는 학습 데이터의 크기의 분포에 따라 straggler 문제가 더욱 심화된다. 이를 확인하기 위하여, 표 1의 학습 데이터 셋 중 UCF101, Charades, DAVIS, HACS의 DIF를 가지면서 학습 데이터의 수와 크기의 평균은 동일하도록 UCF101 데이터셋을 가공하여 DIF가 0일 때에 비해 학습 시간의 증가량을 측정하였다.

도 2는 DIF에 따른 딥러닝 모델별 학습 시간 증가량을 도시한 그래프이다. 도 2를 참조하면, 딥러닝 모델은 VGG16, ResNet50, EfficientNet-B3, MobileNet의 4가지 모델을 사용하였다. 딥러닝 모델의 구조에 따라 학습 시간 증가량의 차이는 존재하였지만, 공통적으로 DIF에 따라 학습 시간이 선형적으로 증가함을 알 수 있다.

이하, 첨부된 도면을 참조하여, 본 발명의 바람직한 실시예에 따른 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 시스템의 구조 및 동작에 대하여 구체적으로 설명한다.

M개의 노드로 구성된 클러스터에서 N개의 데이터로 분산 딥러닝을 수행할 때, i번째 배치에서 j번째 노드의 학습 시간을 $t_{i,j}^{batch}$ 로 정의하면, i번째 배치의 straggler 문제의 크기 SP_i 는 straggler의 학습 시간과 가장 빨리 학습이 완료된 노드의 학습 시간의 차이로 나타낼 수 있으며, 수학식 2와 같이 표현된다.

수학식 2

$$SP_i = \max(t_{i,1}^{batch}, \dots, t_{i,M}^{batch}) - \min(t_{i,1}^{batch}, \dots, t_{i,M}^{batch})$$

[0034] 이때, 도 2의 실험 결과 데이터로부터 총 학습 시간($\sum t_i^{batch}$)은 DIF의 일차식으로 볼 수 있는데, DIF(D)는 D의 표준 편차이므로 총 학습 시간을 수학적 3과 같이 상수 a, b 그리고 $\sum t_i^{batch}$ 에 대한 선형 결합으로 쓸 수 있다.

수학적 3

$$\sum t_i^{batch} = a \sqrt{\frac{1}{|D|-1} \sum_{d \in D} (d - \mu(D))^2} + b$$

[0035]

[0036] 수학적 3에서 우변은 d에 대한 일차식이므로, $t_{i,j}^{batch}$ 를 다시 수학적 4와 같이 i번째 배치에서 j번째 노드의 학습 데이터의 크기 합(λ_{ij})에 대한 선형 결합으로 쓸 수 있다.

수학적 4

$$t_{i,j}^{batch} = a_{ij} \lambda_{ij} + b_{ij}$$

[0037]

[0038] 따라서, 본 발명은 straggler 문제를 최소화하기 위하여 모든 노드들의 학습 시간($t_{i,j}^{batch}$)들을 균일하게 만드는 학습 데이터의 크기 합(λ_{ij})들의 조합을 찾아야 한다. 이를 위하여, 본 발명은 마스트-워커 모델 기반의 데이터 불균형을 고려한 데이터 스케줄러(Imbalance-aware Data Scheduler)를 제안한다.

[0039] 도 3은 본 발명의 바람직한 실시예에 따른 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 시스템을 도시한 구조도이다. 본 발명에 따른 분산 딥러닝 가속화 시스템(1)은 마스터(30) 및 각 노드에 할당된 하나 또는 둘 이상의 워커들(40)을 구비하며, 상기 마스터는 각 노드들의 학습 성능에 따라 학습 데이터를 스케줄링 및 셔플링함으로써 데이터의 불균형을 최소화시키고, 학습률을 보정하여 딥러닝 모델의 수렴 속도를 향상시키게 된다. 상기 마스터 및 워커는 컴퓨팅 시스템에서 실행되는 프로그램으로 구현될 수 있다. 도 3에서, 숫자가 쓰여진 박스는 학습할 데이터를 의미하며, 각 박스에 쓰인 숫자는 해당 데이터의 예상 학습 시간을 나타낸다.

[0040] 각 워커(40)는, 마스터로부터 데이터를 할당받아 노드에서 딥러닝 모델에 대한 학습을 진행하게 된다. 도 3을 참조하면, 워커는 붉은 색으로 표시된 데이터를 할당받고, 이를 이용하여 노드에서 딥러닝 모델을 학습시키게 된다. 각 워커는 할당된 데이터에 대한 딥러닝 모델의 오차를 계산하고 역전파 과정을 통해 기울기(Gradient)를 계산하고, 계산된 기울기를 각 노드의 매개 변수 서버들(42)이 공유하도록 하여 딥러닝 모델의 매개 변수 정보를 동기화하는 과정을 반복하게 된다. 한편, 각 워커는 배치 학습이 종료되면 학습한 데이터의 크기와 학습 시간으로 이루어진 피드백을 마스터로 전송한다.

[0041] 상기 마스터(30)는 노드 성능 평가 모듈(300), 데이터 셔플링 모듈(310), 데이터 스케줄링 모듈(320), 학습률 보정 모듈(도시되지 않음)을 구비하여, 워커에게 각 노드가 학습할 데이터를 할당한다. 마스터는 워커에게 각 노드가 학습할 데이터를 할당하고, 할당된 데이터의 색인 정보를 제공하는데, 초기 학습 단계에서 처음으로 할당하는 데이터는 학습의 무작위성을 위해 임의의 노드에 무작위로 선택된 데이터를 할당하며, 그 후의 학습 단계에서는 각 노드의 예상 학습 시간을 균등하게 분포시킬 수 있도록 하기 위하여 노드 성능 평가, 데이터 셔플링 및 스케줄링하여 각 노드에게 데이터를 할당하게 된다.

[0042] 상기 노드 성능 평가 모듈(300)은, 워커로부터 전송된 학습한 데이터의 크기와 학습 시간의 집합들을 이용하여, 수학적 4로부터 선형 회귀를 통해 노드의 성능 지표 a_{ij} 와 b_{ij} 를 구하고, 상기 성능 지표를 이용하여 각 노드의 성능을 평가한다.

[0043] 상기 데이터 서플링 모듈(310)은, 한 세대의 학습이 완료되면 각 노드가 다음 세대에서 학습할 데이터를 섞는다. 이때, 데이터 이동의 오버헤드를 최소화시키기 위하여 후술되는 데이터 스케줄링 모듈에 사용되는 동일한 알고리즘을 사용하여 최소한의 데이터 서플링만 발생시키는 것이 바람직하다.

[0044] 상기 데이터 스케줄링 모듈(320)은, 데이터의 크기와 노드의 성능 지표를 통해, 각 노드의 데이터 당 예상 학습 시간을 재계산하고 내림차순으로 정렬한 후, 각 노드가 학습할 데이터를 할당하며, 상기 할당된 데이터의 색인 정보를 각 워커에게 제공한다. 워커는 마스터로부터 제공된 데이터 색인 정보를 이용하여, 각 노드에서 다음 배치 학습을 진행하게 된다.

[0045] 상기 데이터 스케줄링 모듈은, 초기 배치 학습을 위하여 처음으로 할당하는 데이터만 학습의 무작위성을 위해 임의의 노드에서 무작위로 데이터를 선택하고, 그 후의 배치 학습에는 각 노드의 예상 학습 시간을 균등하게 분포시키기 위하여 First-fit-decreasing Bin Packing으로 데이터를 선택한다. 모든 노드에서 동일한 시간이 소요되는 데이터 조합을 찾는 알고리즘은 NP-Hard 이므로 본 발명에서는 근사 알고리즘을 사용하여 다항 시간에 노드 당 학습 시간이 최대한 비슷해지는 조합을 찾는다.

[0046] 상기 학습률 보정 모듈은, 각 노드에서 학습한 데이터의 크기에 비례하도록 학습률을 조정하여 학습률 보정을 하게 되며, 학습량에 따른 학습률 보정을 통해 딥러닝 모델의 수렴 속도를 높일 수 있게 된다. 학습률 보정식은 수학적식 5와 같이 나타낸다.

수학적식 5

$$w_{t+1} = w_t - \frac{\eta}{\sum_{j=1}^M \Delta x_{tj}} \sum_{j=1}^M g_{tj} \Lambda_{tj}$$

[0047]

[0048] 여기서, w 는 가중치이며, η 는 학습률(Learning rate)이며, g 는 Gradients이며, M 은 클러스터의 워커들의 개수이다.

[0049] 이하, 도 4를 참조하여, 본 발명에 따른 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 시스템의 동작을 구체적으로 설명한다. 도 4는 본 발명의 바람직한 실시예에 따른 데이터 불균형 최소화를 이용한 분산 딥러닝 가속화 시스템에 있어서, 전체 동작을 순차적으로 도시한 흐름도이다.

[0050] 도 4를 참조하면, 마스터는 초기 학습 단계를 위한 데이터를 무작위로 선택하여 각 노드에 할당하고, 할당된 데이터의 색인 정보들을 각 워커에게 제공하여 노드에서의 딥러닝 모델의 초기 학습을 실행하도록 한다.

[0051] 각 워커는, 마스터로부터 할당된 데이터를 이용하여 노드에서 미니 배치에 대한 학습을 실행하고, 학습이 완료될 때마다, 마스터에게 학습한 데이터의 크기와 학습 시간으로 이루어진 피드백을 전송한다.

[0052] 마스터는, 워커로부터 전송된 학습한 데이터의 크기와 학습 시간의 집합들을 이용하여, 수학적식 4로부터 선형 회귀를 통해 노드의 성능 지표 a_{ij} 와 b_{ij} 를 구한다. 이에 따라, 마스터는 선형 회귀를 통해 구한 성능 지표를 이용하여 각 노드의 성능을 평가한다.

[0053] 다음, 마스터는 모든 노드의 배치 학습과 성능 평가가 끝나면, 수학적식 4에 따라 데이터의 크기와 노드의 성능 지표를 이용하여 각 노드의 데이터 당 예상 학습 시간을 재계산하고 내림차순으로 정렬한다. 그 후 각 노드가 학습할 데이터를 할당하는데, 처음으로 할당하는 데이터만 학습의 무작위성을 위해 임의의 노드에서 무작위로 데이터를 선택하고, 그 후에는 각 노드의 예상 학습 시간을 균등하게 분포시키기 위하여 First-fit-decreasing Bin Packing으로 데이터를 선택한다. 모든 노드에서 동일한 시간이 소요되는 데이터 조합을 찾는 알고리즘은 NP-Hard 이므로 본 발명에서는 근사 알고리즘을 사용하여 다항 시간에 노드 당 학습 시간이 최대한 비슷해지는 조합을 찾는다. 도 4의 데이터 스케줄링 모듈을 참조하면, 각 노드에 할당된 데이터당 예상 학습 시간이 내림차순으로 정렬되며, 각 노드의 예상 학습 시간을 균등하게 분포시키기 위하여, 각 노드의 예상 학습 시간이 24가 되도록 각 노드의 데이터 조합을 설정하게 된다.

[0054] 다음, 마스터는 할당된 데이터의 색인 정보를 각 워커에게 전송한다.

[0055] 각 워커는, 마스터로부터 할당된 데이터의 색인 정보를 각 노드에 전달하여 딥러닝 모델에 대한 다음 배치 학습

을 진행하도록 한다.

- [0056] 마스터는 한 세대의 학습이 완료되면 각 노드가 다음 세대에서 학습할 데이터를 섞는다. 데이터 이동의 오버헤드를 최소화시키기 위하여 데이터 스케줄링할 때와 동일한 알고리즘을 사용하여 최소한의 데이터 서플링만 발생시킨다. 그리고, 데이터를 서플링할 때 다음 세대에서 각 노드의 총 학습 시간이 균일하도록 데이터를 배치한다. 도 4의 데이터 서플링 모듈을 참조하면, 각 노드의 총 학습 시간이 60이 되도록 하기 위하여, 각 노드에서 최소한의 데이터들을 서플링한다.
- [0057] 본 발명에서는 학습 데이터의 크기를 고려하여 데이터의 불균형을 최소화할 뿐 아니라 학습량에 따른 학습률 보정을 통해 딥러닝 모델의 수렴 속도를 높이고자 한다. 학습률 보정은 각 노드에서 학습한 데이터의 크기에 비례하도록 학습률을 조정하여 이루어지며 학습률 보정식은 전술한 수학식 5와 같이 나타낸다.
- [0058] 본 발명에 따른 분산 딥러닝 가속화 시스템의 성능을 평가하기 위하여, 데이터 스케줄러를 TensorFlow v2.6.2상에 구현하였다. 선형 회귀는 Scikit-learn의 선형 회귀 모듈을 사용하였고, 각 데이터 크기에 대한 인터페이스를 추가하기 위하여 Keras의 Sequence 클래스를 재설계하였으며, 학습률 보정 알고리즘은 TensorFlow의 수집 연산 모듈에 구현하였다.
- [0059] 성능 평가를 위하여, 각각 1개의 NVIDIA RTX 2070 SUPER를 갖는 2대의 서버로 구성된 클러스터에서 실험을 진행하였다. 워크 로드는 도 1의 4 가지 딥러닝 모델을 사용하였으며, 각 워크로드별 특징은 다음과 같다. VGG16은 초기 딥러닝 모델로서 학습 시간이 길고 낮은 정확도를 보인다. ResNet50은 Residual Learning을 이용해 gradient vanishing 문제를 해결한 모델로서, 학습 시간은 VGG16보다 짧고 높은 정확도를 보인다. EfficientNet-B3는 한정된 자원으로 최대의 효율을 내기 위하여 고안된 모델로, 학습 시간이 ResNet50보다 짧다. MobileNet은 컴퓨팅 성능이 제한된 환경에서도 잘 동작하도록 연산량을 줄인 모델로서 학습 시간이 가장 짧다. 본 실험에서는 HACS의 DIF를 갖도록 가공한 UCF101 데이터셋을 사용하여 워크로드별 학습 시간 감소 효과와 수렴 속도를 평가하였으며, 도 1의 DIF를 사용하여 RedNet50을 학습시켰을 때 DIF에 따른 학습 시간 감소 효과를 평가하였다.
- [0060] 도 5는 DIF와 워크로드에 따른 학습 시간을 비교한 그래프이다. 학습 시간의 경우 TensorFlow와 데이터 불균형을 고려한 스케줄러를 각각 적용하여 10번의 세대 동안 걸린 학습 시간의 평균으로 계산하였다. DIF가 커질수록 데이터의 불균형을 고려한 스케줄러가 TensorFlow에 비해 상대적으로 성능이 우수하였고 최대 34%까지 학습 시간을 감소시켰다. 또한, DIF가 0인 경우에도 미세한 차이로 TensorFlow에 비해 나은 성능을 보였는데, 이는 각 워크에서 받은 피드백에 기반하여 동적으로 노드의 성능에 맞게 데이터를 할당했기 때문이다.
- [0061] 워크로드에 따른 학습 시간 또한 데이터의 불균형을 고려한 스케줄러가 워크로드에 관계없이 TensorFlow에 비하여 우수한 성능을 보였다. 가장 큰 성능 차이를 보인 워크로드는 EfficientNet-B3로, TensorFlow에 비해 53%의 학습 가속 효과가 있었다. 학습 시간 감소 효과를 설명하기 위해 노드별로 할당된 데이터 간 DIF를 측정해 본 결과, 데이터 불균형을 고려한 데이터 스케줄러에서 DIF의 평균이 최대 72% 감소하였다. 이로써 데이터 불균형을 고려한 스케줄러가 straggler 문제를 최소화하는 데이터 크기의 조합을 잘 찾아내어 학습을 가속화하였다고 볼 수 있다.
- [0062] 도 6은 워크로드에 따른 모델 수렴 속도를 비교한 그래프이다. 모델 수렴 속도의 경우, 10번의 세대 동안 TensorFlow와 수학식 5를 구현한 학습률 보정 알고리즘을 각각 적용하여 10번 세대 동안의 손실 값으로 계산하였다. 이 경우 워크로드에 관계없이 학습률 보정 알고리즘이 더 빠른 수렴 속도를 보였는데, 10번의 세대를 기준으로 가장 가속화된 워크로드는 24%의 가속 효과를 얻은 ResNet50 이었다. 이는 목표 손실까지 도달하는데 필요한 세대의 수를 줄여 주기 때문에 데이터 불균형을 고려한 데이터 스케줄러와 함께 적용할 경우 같은 성능의 모델을 얻기 위한 학습 시간이 현저하게 줄어들 수 있다.
- [0063] 전술한 실험을 통해, 데이터의 불균형한 정도에 따라 학습 시간은 선형적으로 증가하였으며, 이를 해결하기 위하여 본 발명에 따른 데이터 불균형을 고려한 데이터 스케줄링 및 서플링 정책과 학습률 보정 알고리즘을 통해, 워크로드와 데이터 불균형 정도에 관계없이 딥러닝 모델의 성능을 유지하면서 학습을 가속화시킬 수 있게 된다.
- [0064] 이상에서 본 발명에 대하여 그 바람직한 실시예를 중심으로 설명하였으나, 이는 단지 예시일 뿐 본 발명을 한정하는 것이 아니며, 본 발명이 속하는 분야의 통상의 지식을 가진 자라면 본 발명의 본질적인 특성을 벗어나지 않는 범위에서 이상에 예시되지 않은 여러 가지의 변형과 응용이 가능함을 알 수 있을 것이다. 그리고, 이러한 변형과 응용에 관계된 차이점들은 첨부된 청구 범위에서 규정하는 본 발명의 범위에 포함되는 것으로 해석되어

야 할 것이다.

부호의 설명

[0065]

1 : 분산 딥러닝 가속화 시스템

30 : 마스터

40 : 워커

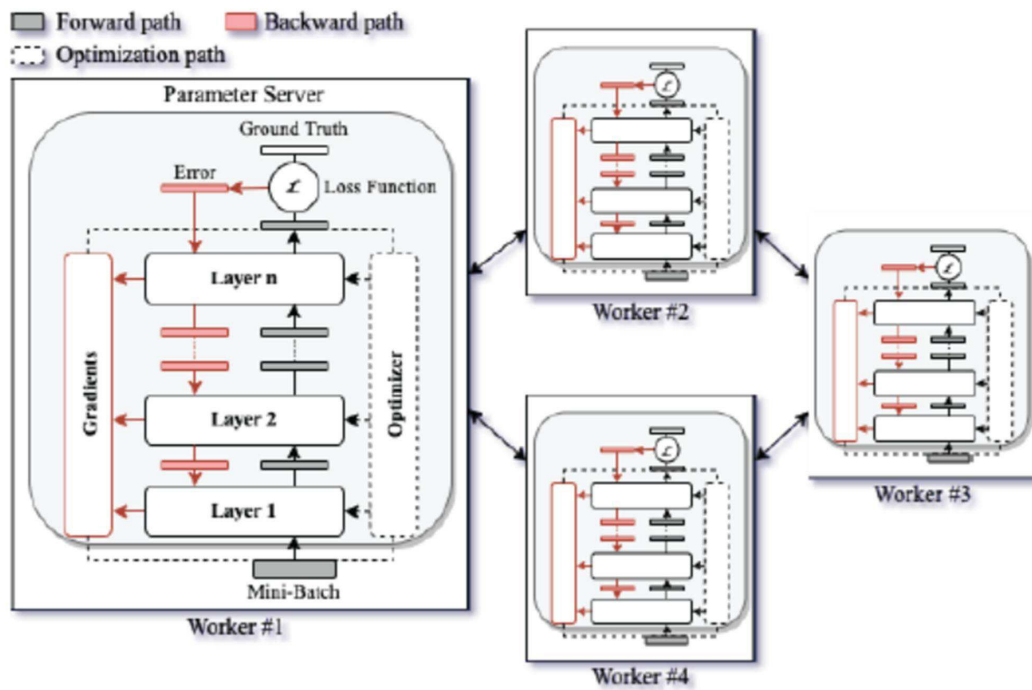
300 : 노드 성능 평가 모듈

310 : 데이터 셔플링 모듈

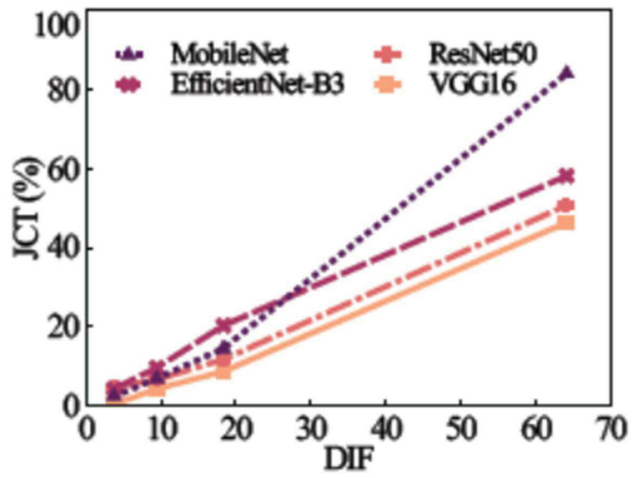
320 : 데이터 스케줄링 모듈

도면

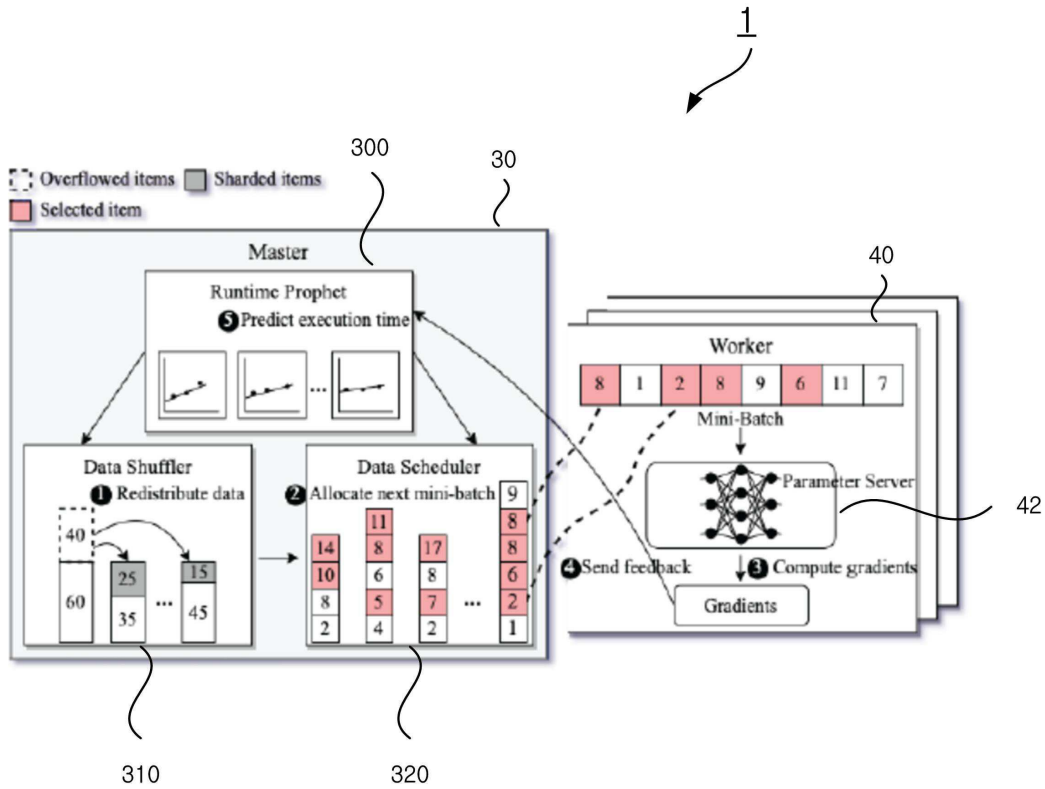
도면1



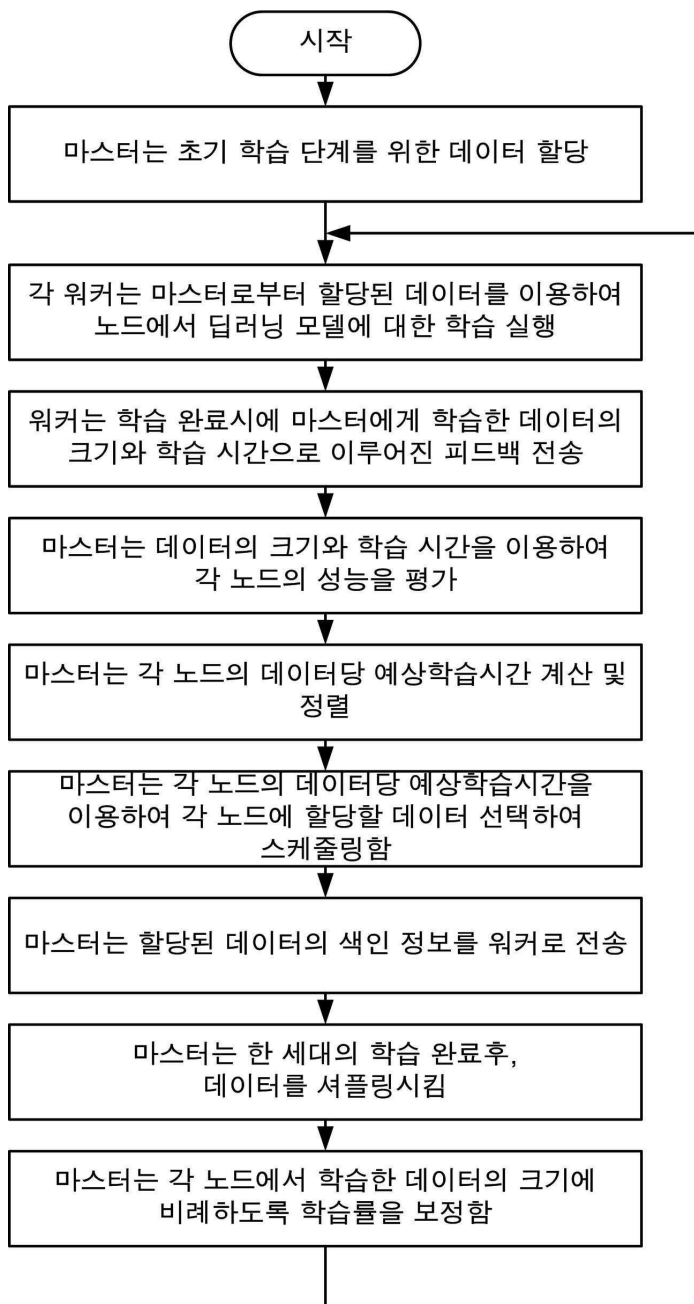
도면2



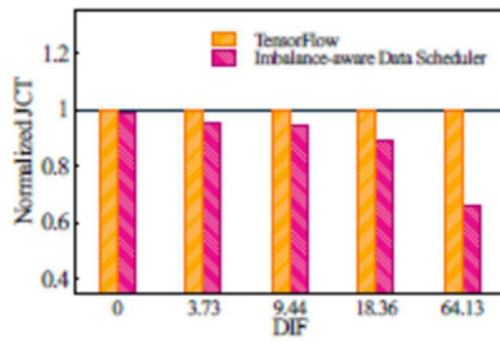
도면3



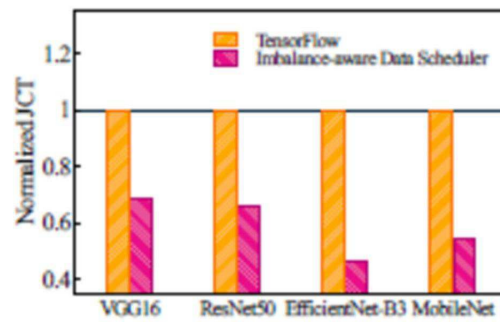
도면4



도면5

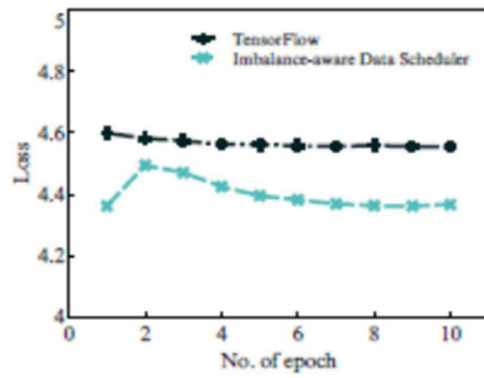


(a) DIF 별 학습 시간.

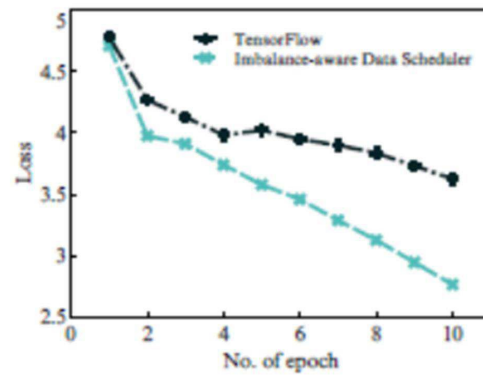


(b) 워크로드 별 학습 시간.

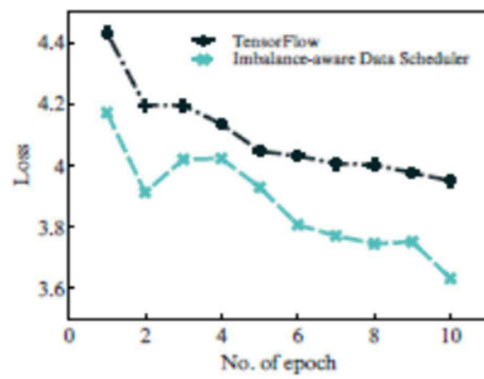
도면6



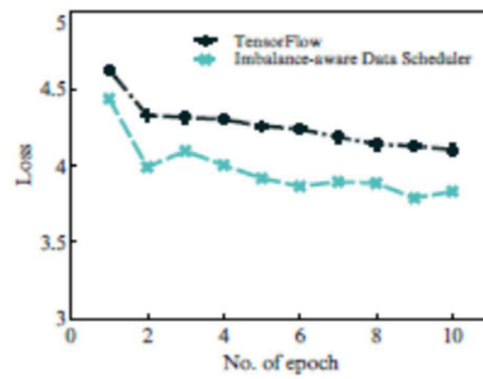
(a) VGG16



(b) ResNet50



(c) EfficientNet-B3



(d) MobileNet