



등록특허 10-2573880



(19) 대한민국특허청(KR)
(12) 등록특허공보(B1)

(45) 공고일자 2023년09월06일
(11) 등록번호 10-2573880
(24) 등록일자 2023년08월29일

(51) 국제특허분류(Int. Cl.)
G06N 3/04 (2023.01) G06N 20/20 (2019.01)
G06N 3/08 (2023.01)
(52) CPC특허분류
G06N 3/045 (2023.01)
G06N 20/20 (2021.08)
(21) 출원번호 10-2022-0090291
(22) 출원일자 2022년07월21일
심사청구일자 2022년07월21일
(56) 선행기술조사문헌
W02021119601 A1*
(뒷면에 계속)

(73) 특허권자
고려대학교 산학협력단
서울특별시 성북구 안암로 145, 고려대학교 (안암
동5가)
(72) 발명자
김중현
서울특별시 동작구 상도로 346-2, 212동201호(상
도동, 상도엠코타운 애스톤파크)
정소이
서울특별시 동대문구 왕산로 86, 101동 1707호(용
두동, 동대문 푸르지오시티)
(뒷면에 계속)
(74) 대리인
윤귀상

전체 청구항 수 : 총 6 항

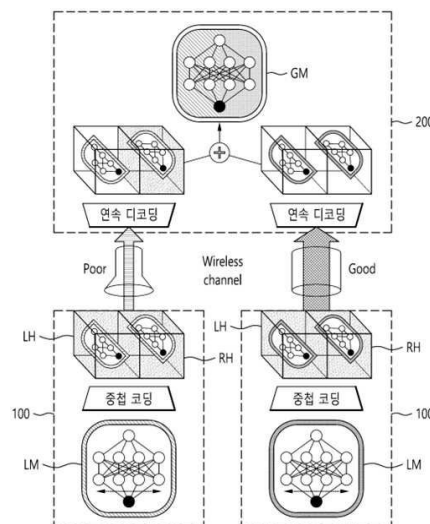
심사관 : 노지명

(54) 발명의 명칭 다중-너비 인공신경망에 기반한 연합 학습 시스템 및 연합 학습 방법

(57) 요약

본 발명은 연합 학습 시스템에 관한 것으로, 기 저장된 제1 로컬 모델로부터 복수의 제2 로컬 모델을 생성하고, 상기 복수의 제2 로컬 모델의 파라미터를 중첩 코딩하여 서버로 전송하는 단말; 및 중첩 코딩된 복수의 제2 로컬 모델의 파라미터를 수신하여 디코딩하고, 디코딩된 파라미터를 집계(aggregates)하여 글로벌 모델을 생성하며, 상기 글로벌 모델의 파라미터를 상기 단말로 전송하는 서버를 포함하고, 상기 단말은, 상기 글로벌 모델의 파라미터에 기초하여 상기 제1 로컬 모델을 업데이트한다. 이에 의해 통신 효율을 높일 수 있는 것은 물론, 데이터 희소성 및 데이터 불균형으로 인해 발생하는 non-IID(non-independent and identically distributed) 문제에 대응할 수 있고, 확인되지 않은 통신 채널 상황에 능동적으로 대처할 수 있다.

대표도 - 도3



(52) CPC특허분류

G06N 3/08 (2023.01)

(72) 발명자

윤원준

서울특별시 강남구 학동로64길 14, 101동 1002호(삼성동, 현대아파트)

백한결

서울특별시 도봉구 도봉로136가길 73, 101동 1304호(창동, 북한산한신희플러스아파트)

곽윤석

서울특별시 서대문구 신촌로9길 37-5, 305호(창천동)

(56) 선행기술조사문헌

J. Hong 등, 'EFFICIENT SPLIT-MIX FEDERATED LEARNING FOR ON-DEMAND AND IN-SITU CUSTOMIZATION', arXiv:2203.09747v1 [cs.LG], 2022.03.18. 1부.*

KR102416342 B1

KR1020210132500 A

KR1020220084508 A

*는 심사관에 의하여 인용된 문헌

이 발명을 지원한 국가연구개발사업

과제고유번호	1711152645
과제번호	2021-0-00467-002
부처명	과학기술정보통신부
과제관리(전문)기관명	정보통신기획평가원
연구사업명	6G핵심기술개발
연구과제명	지능형 6G 무선 액세스 시스템
기 여 율	1/1
과제수행기관명	고려대학교산학협력단
연구기간	2021.04.01 ~ 2025.12.31

공지예외적용 : 있음

명세서

청구범위

청구항 1

기 저장된 제1 로컬 모델로부터 복수의 제2 로컬 모델을 생성하고, 상기 복수의 제2 로컬 모델의 파라미터를 중첩 코딩하여 서버로 전송하는 단말; 및

중첩 코딩된 복수의 제2 로컬 모델의 파라미터를 수신하여 디코딩하고, 디코딩된 파라미터를 집계(aggregates)하여 글로벌 모델을 생성하며, 상기 글로벌 모델의 파라미터를 상기 단말로 전송하는 서버를 포함하고,

상기 단말은,

상기 글로벌 모델의 파라미터에 기초하여 상기 제1 로컬 모델을 업데이트하고, 상기 제1 로컬 모델로부터 너비 변경이 가능한 인공지능 알고리즘을 기초로 상기 복수의 제2 로컬 모델을 생성하되,

상기 복수의 제2 로컬 모델의 너비는 상기 제1 로컬 모델의 너비 이하이며,

상기 복수의 제2 로컬 모델은,

상기 제1 로컬 모델의 절반 너비를 갖는 하프(Half) 제2 로컬 모델 및 상기 제1 로컬 모델의 전체 너비를 갖는 풀(full) 제2 로컬 모델 중 적어도 하나를 포함하고,

상기 하프 제2 로컬 모델은,

상기 제1 로컬 모델의 왼쪽 절반에 대응되는 LH-제2 로컬 모델 및 상기 제1 로컬 모델의 오른쪽 절반에 대응되는 RH-제2 로컬 모델을 포함하는 것을 특징으로 하는 연합 학습 시스템.

청구항 2

삭제

청구항 3

삭제

청구항 4

제1항에 있어서,

상기 서버는,

상기 중첩 코딩된 복수의 제2 로컬 모델의 파라미터를 연속적으로 디코딩하되,

통신 채널의 페이딩(fading) 이득에 기초하는 디코딩 성공 확률에 따라 상기 복수의 제2 로컬 모델의 파라미터에 대한 획득여부를 결정하는 것을 특징으로 하는 연합 학습 시스템.

청구항 5

제4항에 있어서,

상기 서버는,

상기 디코딩 성공 확률에 따라 상기 LH-제2 로컬 모델 및 RH-제2 로컬 모델 중 하나의 파라미터만을 획득하거나 상기 풀 제2 로컬 모델의 파라미터를 획득하는 것을 특징으로 하는 연합 학습 시스템.

청구항 6

제1항에 있어서,

상기 단말은,

상기 글로벌 모델의 파라미터를 기초로 인플레이스 지식 정제(IPKD, inplace knowledge distillation)를 사용하여 상기 제1 로컬 모델을 업데이트하고,

상기 IPKD는,

상기 제1 로컬 모델보다 작은 너비를 갖는 상기 제2 로컬 모델이 상기 제1 로컬 모델과 유사한 로짓(logit)을 산출하도록 하는 방식인 것을 특징으로 하는 연합 학습 시스템.

청구항 7

단말이 기 저장된 제1 로컬 모델로부터 복수의 제2 로컬 모델을 생성하는 단계;

상기 단말이 상기 복수의 제2 로컬 모델의 파라미터를 중첩 코딩하여 서버로 전송하는 단계;

상기 서버가 중첩 코딩된 복수의 제2 로컬 모델의 파라미터를 디코딩하고, 디코딩된 파라미터를 집계(aggregates)하여 글로벌 모델을 생성하는 단계;

상기 서버가 글로벌 모델의 파라미터를 상기 단말로 전송하는 단계; 및

상기 단말이 상기 글로벌 모델의 파라미터에 기초하여 상기 제1 로컬 모델을 업데이트하는 단계를 포함하고,

상기 복수의 제2 로컬 모델을 생성하는 단계에서는,

상기 단말이 상기 제1 로컬 모델로부터 너비 변경이 가능한 인공지능 알고리즘을 기초로 상기 복수의 제2 로컬 모델을 생성하되,

상기 복수의 제2 로컬 모델의 너비는 상기 제1 로컬 모델의 너비 이하이며,

상기 복수의 제2 로컬 모델은,

상기 제1 로컬 모델의 절반 너비를 갖는 하프(Half) 제2 로컬 모델 및 상기 제1 로컬 모델의 전체 너비를 갖는 풀(full) 제2 로컬 모델 중 적어도 하나를 포함하고,

상기 하프 제2 로컬 모델은,

상기 제1 로컬 모델의 왼쪽 절반에 대응되는 LH-제2 로컬 모델 및 상기 제1 로컬 모델의 오른쪽 절반에 대응되는 RH-제2 로컬 모델을 포함하는 것을 특징으로 하는 연합 학습 방법.

청구항 8

삭제

청구항 9

삭제

청구항 10

제7항에 있어서,

상기 디코딩된 파라미터를 집계(aggregates)하여 글로벌 모델을 생성하는 단계에서는,

상기 서버가 상기 중첩 코딩된 복수의 제2 로컬 모델의 파라미터를 연속적으로 디코딩하되,

통신 채널의 페이딩(fading) 이득에 기초하는 디코딩 성공 확률에 따라 상기 복수의 제2 로컬 모델의 파라미터에 대한 획득여부를 결정하는 것을 특징으로 하는 연합 학습 방법.

발명의 설명

기술 분야

본 발명은 다중-너비 인공지능망에 기반한 연합 학습 시스템 및 연합 학습 방법에 관한 것으로, 보다 상세하게는 통신의 효율성을 높일 수 있는 다중-너비 인공지능망에 기반한 연합 학습 시스템 및 연합 학습 방법에 관한

[0001]

것이다.

배경 기술

- [0002] 연합 학습(Federated Learning)은 다수의 단말과 하나의 서버가 협력하여 데이터가 탈중앙화된 상황에서 글로벌 모델(Global model)을 학습하는 기계 학습 기술이다. 여기서, 단말은 예컨대 사물 인터넷 기기, 스마트 폰 등이 될 수 있다.
- [0003] 연합 학습은 제한된 양의 로컬 데이터를 학습한 단말의 로컬 모델을 교환하고 종합하여 학습 샘플의 부족을 극복할 수 있다. 그리고 연합 학습의 효율성을 높이기 위해서는 무선으로 연결되는 기기종 장치인 단말이 많아지도록 해야 한다.
- [0004] 하지만 기기종 장치는 서로 다른 수준의 가용 에너지를 갖게 되므로 상대적으로 가용 에너지가 적은 저 에너지 장치는 너비가 작은 모델을 실행하는 반면, 가용 에너지가 큰 고 에너지 장치는 너비가 큰 모델을 실행할 수 있다.
- [0005] 이에 동일한 아키텍처에서 로컬 모델만 집계할 수 있는 종래의 연합 학습은 동시에 너비가 작은 모델 및 너비가 큰 모델을 학습할 수 없다. 따라서 가용 에너지가 서로 다른 기기종 장치에 대처하기 위해서는 연합 학습을 두 번 수행해야 하며, 이로 인해 전체 학습 시간이 증가한다는 한계가 있다.
- [0006] 한편 연합 학습에서 데이터 희소성 및 데이터 불균형으로 인해 발생하는 non-IID(non-independent and identically distributed) 문제를 줄이기 위해, 학습을 그룹화(grouping, clustering)하는 방법들이 제안되고 있지만, 이 경우에는 학습 샘플이 감소된 복수의 단말 그룹에 대해 연합 학습을 동시에 실행함으로써 정확도가 떨어진다는 문제가 있다.
- [0007] 이러한 문제를 해결하기 위해 인공신경망의 너비 제어가 가능한 SNN(Slimmable Neural Network) 모델을 차용하였다. 이러한 SNN 모델을 사용하여 연합 학습을 수행하면 모든 단말에서 연합하면서 한 번에 훈련할 수 있으며, 그 후에 훈련된 각 로컬 SNN 모델은 너비를 조정할 수 있다.
- [0008] 그리고 통신 채널의 상태는 단말에 따라 다른 것은 물론, 시간에 따라 변경되기 때문에 채널 상태를 알고 있으면, 이러한 SNN 모델을 사용하여 채널 상태가 좋지 않은 단말은 너비가 작은 모델만 교환하고, 채널 상태가 좋은 단말은 너비가 큰 모델을 교환하여 연합 학습에 참여할 수 있다.
- [0009] 따라서 SNN 모델을 사용하면 열악한 채널 상태의 단말은 너비를 줄인 후 로컬 모델을 서버로 보내고, 글로벌 모델 구성의 일부에만 기여하는 방식으로 협력할 수 있다.
- [0010] 하지만 이러한 방법은 임의의 페이딩(fading) 및 이동성으로 인해 시간과 위치에 따라 변하는 채널 상태를 조사하기 위한 추가 통신 및 에너지 비용을 수반한다는 문제가 있다.
- [0011] 따라서 기기종 에너지 용량을 감안할 수 있는 것은 물론, 확인되지 않은 통신 채널 상황에 능동적으로 대처하여 연합 학습을 수행할 수 있는 방법이 필요하다.

선행기술문헌

특허문헌

- [0012] (특허문헌 0001) 한국공개특허공보 제10-2021-0121915호

발명의 내용

해결하려는 과제

- [0013] 본 발명은 상기와 같은 문제를 해결하기 위해 안출된 것으로, 본 발명의 목적은 통신 효율을 높일 수 있는 것은 물론, 데이터 희소성 및 데이터 불균형으로 인해 발생하는 non-IID(non-independent and identically distributed) 문제에 대응할 수 있는 다중-너비 인공신경망에 기반한 연합 학습 시스템 및 연합 학습 방법을 제공하는 것이다.
- [0014] 그리고 본 발명의 다른 목적은 확인되지 않은 통신 채널 상황에 능동적으로 대처할 수 있는 다중-너비 인공신경

망에 기반한 연합 학습 시스템 및 연합 학습 방법을 제공하는 것이다.

과제의 해결 수단

- [0015] 상기 목적을 달성하기 위한 본 발명의 일 실시예에 따른 연합 학습 시스템은, 기 저장된 제1 로컬 모델로부터 복수의 제2 로컬 모델을 생성하고, 상기 복수의 제2 로컬 모델의 파라미터를 중첩 코딩하여 서버로 전송하는 단말; 및 중첩 코딩된 복수의 제2 로컬 모델의 파라미터를 수신하여 디코딩하고, 디코딩된 파라미터를 집계(aggregates)하여 글로벌 모델을 생성하며, 상기 글로벌 모델의 파라미터를 상기 단말로 전송하는 서버를 포함하고, 상기 단말은, 상기 글로벌 모델의 파라미터에 기초하여 상기 제1 로컬 모델을 업데이트한다.
- [0016] 그리고 상기 단말은, 상기 제1 로컬 모델로부터 너비 변경이 가능한 인공지능 알고리즘을 기초로 상기 복수의 제2 로컬 모델을 생성하되, 상기 복수의 제2 로컬 모델의 너비는 상기 제1 로컬 모델의 너비 이하일 수 있다.
- [0017] 또한 상기 복수의 제2 로컬 모델은, 상기 제1 로컬 모델의 절반 너비를 갖는 하프(Half) 제2 로컬 모델 및 상기 제1 로컬 모델의 전체 너비를 갖는 풀(full) 제2 로컬 모델 중 적어도 하나를 포함하고, 상기 하프 제2 로컬 모델은, 상기 제1 로컬 모델의 왼쪽 절반에 대응되는 LH-제2 로컬 모델 및 상기 제1 로컬 모델의 오른쪽 절반에 대응되는 RH-제2 로컬 모델을 포함할 수 있다.
- [0018] 그리고, 상기 서버는, 상기 중첩 코딩된 복수의 제2 로컬 모델의 파라미터를 연속적으로 디코딩하되, 통신 채널의 페이딩(fading) 이득에 기초하는 디코딩 성공 확률에 따라 상기 복수의 제2 로컬 모델의 파라미터에 대한 획득여부를 결정할 수 있다.
- [0019] 또한 상기 서버는, 상기 디코딩 성공 확률에 따라 상기 LH-제2 로컬 모델 및 RH-제2 로컬 모델 중 하나의 파라미터만을 획득하거나 상기 풀 제2 로컬 모델의 파라미터를 획득할 수 있다.
- [0020] 그리고 상기 단말은, 상기 글로벌 모델의 파라미터를 기초로 인플레이스 지식 정제(IPKD, inplace knowledge distillation)를 사용하여 상기 제1 로컬 모델을 업데이트하고, 상기 IPKD는, 상기 제1 로컬 모델보다 작은 너비를 갖는 상기 하프 제2 로컬 모델이 상기 제1 로컬 모델과 유사한 로짓(logit)을 산출하도록 하는 방식일 수 있다.
- [0021] 한편 상기 목적을 달성하기 위한 본 발명의 일 실시예에 따른 연합 학습 방법은, 단말이 기 저장된 제1 로컬 모델로부터 복수의 제2 로컬 모델을 생성하는 단계; 상기 단말이 상기 복수의 제2 로컬 모델의 파라미터를 중첩 코딩하여 서버로 전송하는 단계; 상기 서버가 중첩 코딩된 복수의 제2 로컬 모델의 파라미터를 디코딩하고, 디코딩된 파라미터를 집계(aggregates)하여 글로벌 모델을 생성하는 단계; 상기 서버가 글로벌 모델의 파라미터를 상기 단말로 전송하는 단계; 및 상기 단말이 상기 글로벌 모델의 파라미터에 기초하여 상기 제1 로컬 모델을 업데이트하는 단계를 포함한다.
- [0022] 또한 상기 복수의 제2 로컬 모델을 생성하는 단계에서는, 상기 단말이 상기 제1 로컬 모델로부터 너비 변경이 가능한 인공지능 알고리즘을 기초로 상기 복수의 제2 로컬 모델을 생성하되, 상기 복수의 제2 로컬 모델의 너비는 상기 제1 로컬 모델의 너비 이하일 수 있다.
- [0023] 그리고 상기 복수의 제2 로컬 모델은, 상기 제1 로컬 모델의 절반 너비를 갖는 하프(Half) 제2 로컬 모델 및 상기 제1 로컬 모델의 전체 너비를 갖는 풀(full) 제2 로컬 모델 중 적어도 하나를 포함하고, 상기 하프 제2 로컬 모델은, 상기 제1 로컬 모델의 왼쪽 절반에 대응되는 LH-제2 로컬 모델 및 상기 제1 로컬 모델의 오른쪽 절반에 대응되는 RH-제2 로컬 모델을 포함할 수 있다.
- [0024] 또한 상기 디코딩된 파라미터를 집계(aggregates)하여 글로벌 모델을 생성하는 단계에서는, 상기 서버가 상기 중첩 코딩된 복수의 제2 로컬 모델의 파라미터를 연속적으로 디코딩하되, 통신 채널의 페이딩(fading) 이득에 기초하는 디코딩 성공 확률에 따라 상기 복수의 제2 로컬 모델의 파라미터에 대한 획득여부를 결정할 수 있다.

발명의 효과

- [0025] 상술한 본 발명의 일 측면에 따르면, 다중-너비 인공지능망에 기반한 연합 학습 시스템 및 연합 학습 방법을 제공함으로써, 통신 효율을 높일 수 있는 것은 물론, 데이터 희소성 및 데이터 불균형으로 인해 발생하는 non-IID(non-independent and identically distributed) 문제에 대응할 수 있다. 또한 확인되지 않은 통신 채널 상황에 능동적으로 대처할 수 있다.

도면의 간단한 설명

- [0026] 도 1은 종래의 연합 학습 시스템을 설명하기 위한 도면,
 도 2는 본 발명의 일 실시예에 따른 연합 학습 시스템의 구성을 도시한 도면,
 도 3 및 도 4는 본 발명의 연합 학습 시스템에서 중첩 코딩 및 연속 디코딩을 수행하는 모습을 도시한 도면,
 도 5는 본 발명의 연합 학습 시스템에서 인플레이스 지식 정제(IPKD)가 적용되는 모습을 도시한 도면,
 도 6은 본 발명의 연합 학습 시스템에서 연합 학습을 수행하는 전체적인 과정을 도시한 도면,
 도 7은 도 6의 연합 학습 시스템에서 수행되는 연합 학습 방법을 설명하기 위한 흐름도, 그리고,
 도 8 내지 도 10은 본 발명의 연합 학습 방법의 효과를 설명하기 위한 실험 결과를 도시한 그래프이다.

발명을 실시하기 위한 구체적인 내용

- [0027] 후술하는 본 발명에 대한 상세한 설명은, 본 발명이 실시될 수 있는 특정 실시예를 예시로서 도시하는 첨부 도면을 참조한다. 이들 실시예는 당업자가 본 발명을 실시할 수 있기에 충분하도록 상세히 설명된다. 본 발명의 다양한 실시예는 서로 다르지만 상호 배타적일 필요는 없음이 이해되어야 한다. 예를 들어, 여기에 기재되어 있는 특정 형상, 구조 및 특성은 일 실시예와 관련하여 본 발명의 정신 및 범위를 벗어나지 않으면서 다른 실시예로 구현될 수 있다. 또한, 각각의 개시된 실시예 내의 개별 구성요소의 위치 또는 배치는 본 발명의 정신 및 범위를 벗어나지 않으면서 변경될 수 있음이 이해되어야 한다. 따라서, 후술하는 상세한 설명은 한정적인 의미로서 취하려는 것이 아니며, 본 발명의 범위는, 적절하게 설명된다면, 그 청구항들이 주장하는 것과 균등한 모든 범위와 더불어 첨부된 청구항에 의해서만 한정된다. 도면에서 유사한 참조부호는 여러 측면에 걸쳐서 동일하거나 유사한 기능을 지칭한다.
- [0028] 본 발명에 따른 구성요소들은 물리적인 구분이 아니라 기능적인 구분에 의해서 정의되는 구성요소들로써 각각이 수행하는 기능들에 의해서 정의될 수 있다. 각각의 구성요소들은 하드웨어 또는 각각의 기능을 수행하는 프로그램 코드 및 프로세싱 유닛으로 구현될 수 있을 것이며, 두 개 이상의 구성요소의 기능이 하나의 구성요소에 포함되어 구현될 수도 있을 것이다. 따라서 이하의 실시예에서 구성요소에 부여되는 명칭은 각각의 구성요소를 물리적으로 구분하기 위한 것이 아니라 각각의 구성요소가 수행하는 대표적인 기능을 암시하기 위해서 부여된 것이며, 구성요소의 명칭에 의해서 본 발명의 기술적 사상이 한정되지 않는 것임에 유의하여야 한다.
- [0029] 이하에서는 도면들을 참조하여 본 발명의 바람직한 실시예들을 보다 상세하게 설명하기로 한다.
- [0030] 도 1은 종래의 연합 학습 시스템을 설명하기 위한 도면, 도 2는 본 발명의 일 실시예에 따른 연합 학습 시스템(1)의 구성을 도시한 도면, 도 3 및 도 4는 본 발명의 연합 학습 시스템(1)의 단말(100)과 서버(200)에서 중첩 코딩 및 연속 디코딩을 수행하는 모습을 도시한 도면, 도 5는 본 발명의 연합 학습 시스템(1)의 단말(100)과 서버(200)에서 인플레이스 지식 정제(IPKD)가 적용되는 모습을 도시한 도면, 그리고 도 6은 본 발명의 연합 학습 시스템(1)의 단말(100)과 서버(200)에서 연합 학습을 수행하는 전체적인 과정을 도시한 도면이다.
- [0031] 본 발명의 연합 학습 시스템(1)은, 도 2에 도시된 바와 같이 다수의 단말(100) 및 서버(200)를 포함한다. 그리고 연합 학습 시스템(1)을 구성하는 단말(100) 및 서버(200)는 연합 학습 방법을 수행하기 위한 소프트웨어(어플리케이션)가(이) 설치되어 실행될 수 있다.
- [0032] 단말(100)은 너비 변경이 가능한 다중-너비 인공신경망을 갖는 기기으로써, 예컨대 IoT(Internet of Things) 기기, 서버, 스마트 폰 등을 포함할 수 있다.
- [0033] 그리고 단말(100)은 데이터 세트를 학습하여 제1 로컬 모델(local model)을 생성하고, 저장할 수 있다.
- [0034] 또한 단말(100)은 기 저장된 제1 로컬 모델로부터 일정 비율의 너비를 갖는 복수의 제2 로컬 모델(LM)을 생성할 수 있다. 구체적으로 단말(100)은 학습된 제1 로컬 모델로부터 제1 로컬 모델의 너비 이하인 너비를 갖는 복수의 제2 로컬 모델(LM)을 생성할 수 있다. 예를 들면 1.0x, 0.75x, 0.5x 및 0.25x의 다양한 너비 구성을 갖는 제2 로컬 모델(LM)을 생성할 수 있다.
- [0035] 본 발명에 따른 단말(100)은 1.0x 너비 및 0.5x 너비를 갖는 제2 로컬 모델(LM)을 생성하는 것이 바람직할 수 있으나, 꼭 이에 한정되는 것은 아니다.
- [0036] 단말(100)은 너비 변경이 가능한 인공지능 알고리즘인 다중-너비 인공신경망에 기초하여 복수의 제2 로컬 모델(LM)을 생성할 수 있다.

- [0037] 도 1은 종래의 연합 학습 시스템을 도시한 개략도이고, 도 2는 본 발명의 연합 학습 시스템을 도시한 개략도이다.
- [0038] 종래의 연합 학습 시스템에서는 도 1과 같이 다양한 너비의 구성을 갖는 각각의 로컬 모델(예를 들면 1.0x의 너비를 갖는 로컬 모델(1.0x NN)과 0.5x의 너비를 갖는 로컬 모델(0.5xNN))의 파라미터를 각각 서버(200)로 전송한다. 그리고, 서버(200)는 이를 수신함으로써 0.5x 너비의 글로벌 모델(Global 0.5x NN)과 1.0x 너비의 글로벌 모델(Global 1.0x NN)을 각각 생성한다.
- [0039] 하지만 도 1과 같이 열악한 상태(poor)의 채널에서는 단말(100)에서 전송하는 1.0x 너비 로컬 모델(1.0x NN)의 파라미터가 업링크되지 않는 상황이 발생할 수 있다.
- [0040] 본 발명은 이를 해결하기 위해 안출된 것으로, 본 발명의 연합 학습 시스템(1)은 열악한 상태의 채널에서 파라미터의 일부라도 서버(200)로 전송하여 통신은 물론 학습적인 측면에서 이득을 제공하기 위해 마련된다.
- [0041] 단말(100)은 제1 로컬 모델 너비 이하의 너비를 갖는 복수의 제2 로컬 모델(LM)을 생성할 수 있다.
- [0042] 구체적으로 본 발명에서는 0.5x 및 1.0x 두 가지 너비를 갖는 제2 로컬 모델(LM)을 생성하는 것이 바람직하나, 꼭 이에 한정되는 것은 아니다.
- [0043] 이러한 복수의 제2 로컬 모델(LM)은, 제1 로컬 모델의 절반 너비를 갖는 하프(Half) 제2 로컬 모델 및 제1 로컬 모델의 전체 너비를 갖는 풀(full) 제2 로컬 모델 중 적어도 하나를 포함할 수 있다.
- [0044] 그리고 하프 제2 로컬 모델은, 제1 로컬 모델의 전체 너비에 대해 왼쪽의 절반 너비를 갖는 LH(Left Half)-제2 로컬 모델(LH) 및 상기 제1 로컬 모델의 전체 너비에 대해 오른쪽의 절반 너비를 갖는 RH(Right Half)-제2 로컬 모델(RH)을 포함할 수 있다.
- [0045] 그리고 단말(100)은 복수의 제2 로컬 모델(LM)의 파라미터를 중첩 코딩한 다음, 중첩 코딩된 복수의 제2 로컬 모델(SC-LM)의 파라미터에 다른 전송 전력 수준을 할당하여 서버(200)로 전송되도록 할 수 있다.
- [0046] 또한 단말(100)은 연합 학습을 위해 서버(200)에서 생성된 글로벌 모델(GM)의 파라미터를 서버(200)로부터 수신할 수 있다. 이에 단말(100)은 수신한 글로벌 모델(GM)의 파라미터에 기초하여 제1 로컬 모델을 업데이트할 수 있다.
- [0047] 단말(100)은 이러한 제1 로컬 모델의 업데이트를 연합 학습이 완료될 때까지 반복할 수 있다.
- [0048] 한편 서버(200)는 중첩 코딩된 복수의 제2 로컬 모델(SC-LM)의 파라미터를 수신한다. 그리고 서버(200)는 수신된 복수의 제2 로컬 모델(SC-LM)의 파라미터를 디코딩할 수 있다.
- [0049] 이때 서버(200)는 RH-제2 로컬 모델(RH)의 파라미터에 대한 디코딩을 시도한다. 그리고 서버(200)는 RH-제2 로컬 모델(LH)의 파라미터에 대한 디코딩이 성공하면, 연속하여 LH-제2 로컬 모델(LH)의 파라미터에 대한 디코딩을 시도할 수 있다. 이러한 디코딩은, 본 발명에서 연속 디코딩(Successive Decoding) 또는 연속 간섭 취소(SIC, Successive Interference Cancellation)로 설명될 수 있다.
- [0050] 따라서 단말(100)-서버(200) 간의 채널 처리량이 낮을 때 서버(200)는 단말(100)로부터 수신되는 중첩 코딩된 복수의 제2 로컬 모델(SC-LM)의 파라미터 중에서 RH-제2 로컬 모델(RH)의 파라미터만 디코딩하여 RH-제2 로컬 모델(RH)을 얻을 수 있게 된다.
- [0051] 한편 채널 처리량이 높을 때 서버(200)는 RH-제2 로컬 모델(RH)의 파라미터와 LH-제2 로컬 모델(LH)의 파라미터를 모두 디코딩하고 이들을 결합하여 풀 제2 로컬 모델을 얻을 수 있게 된다.
- [0052] 결과적으로 서버(200)는 연합 학습 시스템(1)을 구성하는 복수의 단말(100)들로부터 디코딩된 하프 제2 로컬 모델의 파라미터와 풀 제2 로컬 모델의 파라미터를 집계하여 글로벌 모델(GM)을 생성할 수 있다.
- [0053] 그리고 서버(200)는 글로벌 모델(GM)의 파라미터를 단말(100)로 전송할 수 있다.
- [0054] 이하에서는 이러한 과정을 보다 구체적으로 설명하기로 한다.
- [0055] 단말(100)은 도 4에 도시된 바와 같이 복수의 제2 로컬 모델(LM) 중 LH-제2 로컬 모델(LH)과 RH-제2 로컬 모델(RH)의 파라미터를 중첩 코딩(Superposition Coding)하고, 중첩 코딩된 복수의 제2 로컬 모델(SC-LM)의 파라미터를 서로 다른 전송 전력으로 서버(200)로 전송할 수 있다. 다시 말해서 단말(100)은 LH-제2 로컬 모델(LH)의 파라미터인 LH 세그먼트와 RH-제2 로컬 모델(RH)의 파라미터인 RH 세그먼트를 중첩 코딩하고, 중첩 코딩된 제2

로컬 모델(SC-LM)의 파라미터는 서버(200)로 전송될 수 있다.

[0056] 그리고 단말(100)은 서버(200)로부터 글로벌 모델(GM)의 파라미터를 수신하면, 글로벌 모델(GM)의 파라미터에 기초하여 제1 로컬 모델을 업데이트함으로써 연합 학습을 수행할 수 있다.

[0057] 또한 단말(100)에서 생성되는 복수의 제2 로컬 모델(LM) 중 적어도 하나는 제1 로컬 모델의 너비와는 다른 너비를 갖도록 구성되므로 이에 기초하여 학습하는 경우 중첩되는 너비로 인해 서로 간섭할 수 있다. 또한 데이터 희소성 및 데이터 불균형으로 인해 발생하는 non-IID(non-independent and identically distributed) 문제가 있다.

[0058] 이에 단말(100)은 서로 다른 너비를 갖는 복수의 제2 로컬 모델(LM) 간의 불필요한 간섭을 피함으로써 데이터 분포에 관계없이 높은 정확도로 빠른 수렴을 달성하는 학습 방법을 수행할 수 있다.

[0059] 구체적으로 t번째 반복(iteration)에서 단말(100)은 두 가지 너비 구성을 가진 가중치 벡터 θ_t^k 를 갖게 된다. 두 가지 너비 구성은 0.5x 너비 구성인 $\theta_t^k \odot \Xi_1$ 과 1.0x 너비 구성인 $\theta_t^k \odot \Xi_2 = \theta_t^k$ 이고, 여기서 $\odot$ 은 원소별 곱(elementwise product)이고, Ξ_i 는 i번째 너비 구성의 파라미터를 추출하기 위한 이진 마스크를 의미할 수 있다.

[0060] 그리고 이전에 훈련되어 기 저장된 제1 로컬 모델의 너비 가중치와는 다른 너비 가중치를 갖는 중첩되는 너비가, 후자의 역전파(BP, backpropagation)에 의해 왜곡될 수 있기 때문에 다중-너비 인공신경망을 학습하는 것은 어렵다. 이러한 너비 간의 간섭은 추론 정확도를 저하시킬 뿐만 아니라, 학습 수렴을 방해한다.

[0061] 너비 변경이 가능한 인공지능 알고리즘에 기초하여 복수의 제2 로컬 모델(LM)을 생성하는 본 발명의 단말(100)은 이를 해결하기 위해 도 5와 같이 제1 로컬 모델에서 하프 제2 로컬 모델(LM)에 대한 인플레이스 지식 정제(IPKD, inplace knowledge distillation)를 추가로 적용한다.

[0062] 인플레이스 지식 정제는 하위 너비(1.0x의 전체 너비보다 작은 너비)의 구성을 갖는 로컬 모델이 전체 너비의 구성을 갖는 로컬 모델의 출력과 유사한 소프트맥스 출력(즉, 로짓(logit))을 산출하도록 하여 중첩되는 역전파 기울기(gradient)의 차이가 감소하여 너비 간 간섭이 감소하도록 하는 방법이다.

[0063] 이에 본 발명의 단말(100)은 non-IID 데이터 분포 문제를 해결하기 위해 먼저 모든 전방 전파(FP, forward propagation) 손실을 유지한 다음, 중첩된 기울기로 제1 로컬 모델을 동시에 업데이트 한다.

[0064] 이를 통해 하프 제2 로컬 모델(LM)은 제1 로컬 모델의 로짓과 불일치하는 로짓이 없도록 인플레이스 지식 정제를 사용하여 학습되는 반면, 제1 로컬 모델은 실측 자료(ground truth)를 사용하여 학습된다.

[0065] 이에 복수의 단말(100)에서의 제1 로컬 모델의 업데이트 규칙은 하기의 수학적 식 1과 같이 표현될 수 있다.

[0066] [수학적 식 1]

$$\theta_{t+1}^k = \theta_t^k - \eta_t [w_1 \nabla \hat{F}^k(\theta_t^k \odot \Xi_1, \zeta_t^k) + w_2 \nabla F^k(\theta_t^k \odot \Xi_2, \zeta_t^k)]$$

[0067] 여기서 w_i 는 i번째로 가장 작은 너비를 갖는 로컬 모델을 통해 제1 로컬 모델을 업데이트하기 위한 상수로, w_1 및 w_2 는 $w_1, w_2 > 0$ 인 상수이고, $w_1 + w_2 = 1$ 이다. 그리고 η_t 는 학습률로 $\eta_t > 0$ 이고, ζ_t^k 는 확률적 입력 실현을 의미한다. 함수 $F^k(\theta_t^k \odot \Xi_i, \zeta_t^k)$ 는 i번째 너비 구성의 지상 실측(ground truth) $y(\zeta_t^k)$ 과 로짓 $M(\theta_t^k \odot \Xi_i, \zeta_t^k)$ 간의 교차 엔트로피(cross-entropy)이다. 한편 인플레이스 지식 정제 함수 $\hat{F}^k(\theta_t^k \odot \Xi_1, \zeta_t^k)$ 는 전체 너비 구성의 로짓 $M(\theta_t^k, \zeta_t^k)$ 간의 교차 엔트로피이다.

[0069] 그리고 본 발명의 단말(100)에서 제1 로컬 모델을 생성하기 위한 구체적인 학습 알고리즘은 하기의 알고리즘 1과 같다.

Algorithm 1: Superposition Training (SUSTrain)

```

1 Initialize train parameter  $\Theta = \{\theta^1, \dots, \theta^k, \dots, \theta^K, \theta^G\}$ .
2 Initialize local dataset  $Z = \{Z_1, \dots, Z_k, \dots, Z_K\}$  with Dirichlet
  distribution
3 Initialize learning rate  $\eta_t \leftarrow \eta_0$ 
4 Further Constraints:  $\hat{F}^1(\cdot) = F^1(\cdot)$  ▷ Discuss in Sec. V
5 for  $t = 1, \dots, T$  do
6   for  $k = 1, \dots, K$  do
7     Initialize gradients of model optimizer as 0.
8     Sample batch  $\zeta_t^k$  from  $Z_k$ .
9     Compute loss,  $loss \leftarrow F^k(\theta_t^k \odot \Xi_i, \zeta_t^k)$ .
10    Execute full-network  $M(\theta_t^k, \zeta_t^k)$ .
11    Accumulate gradients,  $loss.backward()$ .
12    Execute full-network  $M(\theta_t^k, \zeta_t^k)$ .
13    Stop gradients of  $M(\theta_t^k, \zeta_t^k)$  as label.
14    for  $i = 1, \dots, S - 1$  do
15      Execute and calculate loss  $\hat{F}^k(\theta_t^k \odot \Xi_i, \zeta_t^k)$ 
16       $loss \leftarrow loss + w_i \hat{F}^k(\theta_t^k \odot \Xi_i, \zeta_t^k)$ .
17    end
18    Calculate gradient of  $loss$ .
19    Update model parameters. ▷ Eq. (1)
20  end
21 end

```

[0070]

[0071]

한편 서버(200)는 다수의 단말(100)로부터 복수의 제2 로컬 모델(LM)의 파라미터를 전달받아 제2 로컬 모델(LM)의 파라미터를 집계하고, 글로벌 모델(GM)을 생성할 수 있다. 또한 서버(200)는 글로벌 모델(GM)의 파라미터를 단말로 전송할 수 있다.

[0072]

그리고 단말(100)로부터 전달받는 복수의 제2 로컬 모델(LM)의 파라미터는 중첩 코딩된 파라미터이다. 이에 서버(200)는 중첩 코딩된 제2 로컬 모델(SC-LM)의 파라미터를 채널 처리량에 기초하여 연속 디코딩을 수행함으로써 중첩 코딩된 제2 로컬 모델(SC-LM)의 파라미터를 디코딩할 수 있다.

[0073]

구체적으로 서버(200)는 단말(100)로부터 수신한 로컬 모델의 파라미터에 기초하여 디코딩을 시도할 수 있다. 이에 서버(200)는 먼저 LH 세그먼트에 대한 디코딩이 성공하면, 연속적으로 RH 세그먼트에 대한 디코딩을 시도하는 연속 디코딩을 수행할 수 있다.

[0074]

따라서 연합 학습을 수행함에 있어 단말(100)과 서버(200) 간의 채널 처리량이 낮을 때 서버(200)는 LH-제2 로컬 모델(LH)의 파라미터인 LH 세그먼트만 디코딩하여 0.5x 너비의 LH-제2 로컬 모델(LH)을 얻을 수 있다. 한편 채널 처리량이 높을 때는 서버(200)는 LH-제2 로컬 모델(LH)의 파라미터인 LH 세그먼트 및 RH-제2 로컬 모델(RH)의 파라미터인 RH 세그먼트 모두를 디코딩하여 이를 결합함으로써 제1 로컬 모델과 동일한 1.0x의 너비를 갖는 풀 제2 로컬 모델을 얻을 수 있다.

[0075]

따라서 서버(200)는 디코딩된 하프 제2 로컬 모델과 디코딩된 제2 로컬 모델을 중첩하는 글로벌 모델(GM)을 생성할 수 있다. 이후 서버(200)는 글로벌 모델(GM)의 파라미터를 단말(100)로 전달하게 되며, 단말(100)은 서버(200)로부터 다운로드한 글로벌 모델(GM)로 제1 로컬 모델을 업데이트할 수 있다. 이러한 과정은 연합 학습이 완료될 때까지 반복하게 된다.

[0076]

그리고 이상의 단말(100)과 서버(200)를 포함하는 연합 학습 환경에서 글로벌 모델(GM)은 서로 다른 너비를 갖는 모델들이 혼합되어 있기 때문에 표준 연합 학습 수렴이 의심된다. 또한 단말(100)의 복수의 제2 로컬 모델(LM)은 여러 너비로 구성되므로 이를 학습할 때 서로 간섭할 수 있다는 문제가 있다.

[0077]

이에 서버(200)에서의 신호 대 간섭 플러스 잡음비인 SINR(signal-to-interference-plus-noise ratio)은 $\gamma = gd^{-\beta}P/(\sigma^2 + P^I)$ 로 지정된다. 여기서 P, P^I 및 σ^2 는 전송 간섭, 수신 간섭, 그리고 잡음세기(noise powers)를 나타낸다.

[0078]

그리고 d는 서버(200)와 단말(100) 간의 거리, β 는 경로 손실 지수로 $\beta \geq 2$ 이고, g는 임의의 소규모 페이딩(random small-scale fading) 전력 이득이다. 가우시안 코드북(Gaussian codebook)이 있는 샤넬의 용량 공식(Shannon's capacity formula)에 따라 대역폭 W가 있는 수신 처리량 R은

$R = W \log_2(1 + \gamma)$ (bits/sec) 로 정의될 수 있다.

[0079] 송신기인 단말(100)이 코드 비율(code rate) u 로 로(raw) 데이터인 복수의 제2 로컬 모델(LM)의 파라미터를 인코딩할 때 수신기인 서버(200)는 $R > u$ 인 경우 인코딩된 데이터를 성공적으로 디코딩할 수 있다. 이러한 디코딩 성공 확률은 하기의 수학식 2로 표현될 수 있다.

[0080] [수학식 2]

[0081]

$$\Pr(R \geq u) = \Pr\left(\frac{gd^{-\beta}P}{\sigma^2 + P^I} \geq u'\right)$$

[0082]

여기서 u' 는 $u' = 2^{\frac{u}{W}} - 1$ 이다.

[0083]

한편 중첩 코딩과 연속 디코딩을 사용한 디코딩 성공 확률은 후술하는 바와 같이 P 와 P 를 뺄런싱하여 주어진다.

[0084]

본 발명에서는 단말(100)에서 서버(200)로 S 개의 메시지, 즉 복수의 제2 로컬 모델(LM)의 파라미터를 동시에 전송하는 것을 고려한다.

[0085]

이러한 S 개의 메시지는 서버(200)로 전송되기 전에 중첩 코딩되는 반면, 총 전송 전력 예산 P 는 $i \in [1, S]$ 에 대한 $\bar{P} = \sum_{i=1}^S P_i$ 전송 전력의 양으로 i 번째 메시지에 할당된다. $S=1$ 인 경우, 즉 단일 메시지만 전송할 때에는 $P^I = 0$ 이므로, 수신 시 간섭이 존재하지 않는다. 그리고 $S>1$ 인 경우에는 연속 디코딩은 간섭을 결정한다.

[0086]

서버(200)는 단말(100)로부터 수신한 중첩 코딩된 메시지 중에서, 먼저 가장 강한 신호를 디코딩한 다음 재구성된 신호, 즉 가장 강한 신호를 제거하고, 차순으로 강한 신호를 디코딩하여 연속적으로 재구성할 수 있다.

[0087]

송신기인 단말(100)과 수신기인 서버(200) 사이에 LOS(Line-of-Sight)가 없고 수신 경로가 다중의 반사 산란 경로만 존재하는 경우에 나타나는 페이딩인 레일리(Rayleigh) 페이딩에서 소규모 페이딩 전력 이득인 g 는 지수 분포, 즉 $g \sim \text{Exp}(1)$ 를 따른다.

[0088]

구체적으로 $P_i > P_{i'}$ 이고 $i' > i$ 라고 가정하면, 서버(200)는 나머지 메시지를 간섭 P_i^I 로 경험하면서 i 번째 메시지를 연속적으로 디코딩할 수 있다. 이때 $P_i^I = gd^{-\beta} \hat{P}_i^I$ 이고, 여기서 $\hat{P}_i^I \triangleq \sum_{i'=i+1}^S P_{i'}$ 으로 $i \leq S-1$ 이며, 마지막 메시지에 대한 간섭이 없으므로 $\hat{P}_S^I = P_S^I = 0$ 이다. 그리고 i 번째 메시지에 대한 수신 처리량을 R_i 라고 하고, P_i^I 를 상술한 수학식 2에 대입하면 R_i 의 분포는 $P_r(R \geq u) = P_r(g \geq \frac{c}{P_i/u' - \hat{P}_i^I})$ 로 변환되며, 여기서 $c = \sigma^2 d^\beta$ 이다.

[0089]

이러한 결과를 적용하면 i 번째 메시지의 디코딩 성공 확률 p_i 는 수학식 3과 같을 수 있다.

[0090]

[수학식 3]

[0091]

$$\begin{aligned} p_i &= \Pr(R_1 \geq u, R_2 \geq u, \dots, R_i \geq u) \\ &= \Pr\left(g \geq \frac{c}{P_1/u' - \hat{P}_1^I}, \dots, g \geq \frac{c}{P_i/u' - \hat{P}_i^I}\right) \\ &= \Pr\left(g \geq \max\left\{\frac{c}{P_1/u' - \hat{P}_1^I}, \dots, \frac{c}{P_i/u' - \hat{P}_i^I}\right\}\right) \end{aligned}$$

[0092]

즉, 첫번째 메시지가 디코딩된 후 순차적으로 디코딩되며, 본 발명에 따른 메시지의 개수는 2개뿐이므로 i 는 2가 된다.

[0093]

그리고 단말(100) 및 서버(200)를 포함하는 연합 학습 시스템(1)은 하기의 알고리즘 2를 통해 도 6에서와 같이

연합 학습을 수행할 수 있다.

Algorithm 2: SlimFL with SC & SD

```

1 Initialize train parameters  $\Theta = \{\theta^1, \dots, \theta^k, \dots, \theta^K, \theta^G\}$ .
2 Split dataset  $Z$  into  $K$  datasets  $Z = \{Z_1, \dots, Z_k, \dots, Z_K\}$ .
3 while Training do
4   //Local Model Training (Algorithm 1)
5   for  $n = 1, \dots, K$  do
6     for  $\zeta_k$  in  $Z_k$  do
7       Update local model parameter  $\theta^k$  ▷ Eq. (1)
8     end
9   end
10  //SC&SD-based Server Aggregation (Uplink)
11  if Aggregation Period then
12     $n_L = |H \cup F| \leftarrow 0, n_R = |F| \leftarrow 0, H \leftarrow \emptyset, F \leftarrow \emptyset$ 
13    for  $k = 1, \dots, K$  do
14       $g \leftarrow rand(1)$ 
15      if  $p_2 \leq g < p_1$  then
16         $H \leftarrow H \cup k, n_L \leftarrow n_L + 1$ 
17      end
18      if  $g \geq p_2$  then
19         $F \leftarrow F \cup k, n_L \leftarrow n_L + 1, n_R \leftarrow n_R + 1$ 
20      end
21    end
22    ▷ Case1.  $n_L > 0, n_R > 0$ 
23     $\theta^G \leftarrow \frac{1}{|H \cup F|} \sum_{k \in H \cup F} \theta^k \odot \Xi + \frac{1}{|F|} \sum_{k \in F} \theta^k \odot \Xi^{-1}$ 
24    ▷ Case2.  $n_L > 0, n_R = 0$ 
25     $\theta^G \leftarrow \frac{1}{n_L} \sum_{k \in H} (\theta^k \odot \Xi)$ 
26    ▷ Case3.  $n_L = n_R = 0$ 
27    Skip aggregation
28  end
29  //Local Update (Downlink)
30  for  $n = 1, \dots, K$  do
31     $\theta^k \leftarrow \theta^G$ 
32  end
33 end

```

구체적으로, 단말(100)은 K개로 마련되고, 서버(200)는 K개의 단말(100)과 연결될 수 있다. 이에 따라 $k \in [1, K]$ 인 k번째 단말(100)은 로컬 데이터 세트 $Z^k \in Z$ 와 2개의 너비 구성을 갖는 제2 로컬 모델(LM)의 파라미터 θ^k 를 갖는다.

서버(200)가 복수의 단말(100)로부터 수신한 제2 로컬 모델(LM)의 파라미터를 집계한 글로벌 데이터 Z는 단말(100) 전반에 걸쳐 IID(independent identically distributed) 이거나 non-IID일 수 있다.

그리고 파라미터 θ^k 는 LH 세그먼트 $\theta^k \odot \Xi$ 와 RH 세그먼트 $\theta^k \odot \Xi^{-1}$ 로 나뉠 수 있다. 여기서, $\Xi = \Xi_1$ 이고, $\Xi^{-1} = \Xi_2 - \Xi_1$ 이다.

k번째 단말(100)은 알고리즘 2의 4행 내지 9행에 기재된 바와 같이 상술한 알고리즘 1 및 수학식 1에 기초하여 제1 로컬 모델의 파라미터를 업데이트할 수 있다.

이후 단말(100)은 제2 로컬 모델의 파라미터인 θ^k 를 중첩 코딩한 후 서버(200)로 전송할 수 있다.

이때 단말(100)은 $P_1 \gg P_2$ 로써 서로 다른 전송 전력인 P_1 과 P_2 로 2개의 메시지인 LH 세그먼트 및 RH 세그먼트를 전송한다.

이후 서버(200)는 LH 세그먼트 및 RH 세그먼트를 포함하는 중첩 코딩된 파라미터에 대해 수학식 3에 따라 연속 디코딩을 수행함으로써 파라미터를 집계하여 글로벌 모델(GM)을 생성할 수 있다.

그리고 서버(200)에서의 제2 로컬 모델(LM)의 파라미터에 대한 획득 여부는 알고리즘 2에서 15행 내지 17행에 기재된 바와 같이, 통신 채널의 페이딩 이득에 기초해 산출되는 디코딩 성공 확률에 따라 결정될 수 있다.

[0103] 구체적으로 서버(200)는 페이딩 이득 $g \geq c/(P_1/u' - P_2)$ 인 경우에는 너비가 0.5x인 하프 제2 로컬 모델을 획득할 수 있고, 페이딩 이득 $g \geq \max\{c/(P_1/u' - P_2), c/(P_2/u')\}$ 을 만족하는 경우에는 18행 내지 20행에 기재된 바와 같이 너비가 1.0x인 풀 제2 로컬 모델을 획득할 수 있다. 상술한 두 가지의 경우가 아니면 서버(200)는 제2 로컬 모델(LM)을 획득하지 못하게 된다.

[0104] 따라서 서버(200)는 복수의 단말(100)들 중에서 도 6과 같이 성공적으로 디코딩된 RH 세그먼트 세트인 F에서 RH 세그먼트를 집계하고, 성공적으로 디코딩된 LH 세그먼트 세트인 H와 F의 HUF에서 LH 세그먼트를 집계할 수 있다.

[0105] 이때 $|HUF| \approx Kp_1$ 및 $|F| \approx Kp_2$ 가 될 정도로 단말(100)의 개수 K가 충분히 크다고 가정한다. 여기서 p_1 과 p_2 는 각각 수학적 3에서 제시된 LH 세그먼트와 RH 세그먼트의 디코딩 성공 확률이다.

[0106] 결과적으로 t번째 통신 라운드(round)에서 서버(200)는 하기의 수학적 4와 같이 글로벌 모델 θ_t^G 를 생성할 수 있다.

[0107] [수학적 4]

$$\theta_t^G \leftarrow \frac{1}{Kp_1} \sum_{k \in HUF} \theta_t^k \odot \Xi + \frac{1}{Kp_2} \sum_{k \in F} \theta_t^k \odot \Xi^{-1}$$

[0109] 이상의 과정을 수행하는 본 발명의 연합 학습 시스템(1)은 다양한 학습 및 통신 기술을 통합할 수 있을 만큼 충분히 유연하지만, 이후로는 다음의 가정을 고려하여 범위를 제한할 수 있다.

[0110] 먼저 알고리즘 2에서 29행 내지 32행에 기재된 다운로드 디코딩은 중첩 코딩 및 연속 디코딩을 무시하고 항상 성공할 수 있다. 이는 서버(200)가 업링크 전력보다 훨씬 더 큰 전송 전력을 갖는다는 사실에 근거할 수 있다.

[0111] 한편 통신 라운드당 로컬 반복 횟수는 1이므로 위 첨자 G를 생략할 수 있다. 즉 $\theta_t = \theta_t^G$ 이다.

[0112] 이를 통해 본 발명의 연합 학습 시스템(1)은 통신 효율을 높일 수 있는 것은 물론, 데이터 희소성 및 데이터 불균형으로 인해 발생하는 non-IID 문제에 대응할 수 있다. 또한 확인되지 않은 통신 채널 상황에 능동적으로 대처할 수 있다.

[0113] 그리고 상술한 수학적식들에 대한 기호는 하기의 표 1과 같을 수 있다.

[0114] [표 1]

Notation	Description
T	Total iteration steps.
S	The number of SNN width configurations.
θ^G	Global model parameter vector.
θ^k	Local model parameter vector of the k -th device.
$\mathcal{K}$	A set of devices ($\mathcal{K} = \{1, \dots, k, \dots, K\}$).
Ξ	A binary mask to extract weight parameters of an LH segment.
Ξ^{-1}	A binary mask to extract weight parameters of an RH segment.
H	A set of successfully decoded LH segments.
F	A set of successfully decoded RH segments.
n_L	The number of successfully decoded LH segments.
n_R	The number of successfully decoded RH segments.
p_1	The decoding success probability of an LH segment.
p_2	The decoding success probability of an RH segment.
Z	Entire dataset.
Ξ_i	A binary mask to extract weight parameters of the i -th smallest model.
ω_i	Positive constant for updating full model via the i -th smallest model.
η_t	Learning rate at the iteration t .
ζ_t^k	The local data sampled from k -th user at the iteration t .

[0115]

[0117] 한편 도 7은 본 발명의 일 실시예에 따른 연합 학습 방법을 설명하기 위한 흐름도로써, 본 발명의 일 실시예에

따른 연합 학습 방법은 도 2 내지 도 6에 도시된 연합 학습 시스템(1)과 실질적으로 동일한 구성 상에서 진행되므로, 도 2내지 도 6의 연합 학습 시스템(1)과 동일한 구성요소에 대해 동일한 도면 부호를 부여하고, 반복되는 설명은 생략하기로 한다.

- [0118] 본 발명의 연합 학습 방법은, 단말(100)이 제1 로컬 모델로부터 복수의 제2 로컬 모델(LM)을 생성하는 단계(S110), 단말(100)이 복수의 제2 로컬 모델의 파라미터를 중첩 코딩하여 서버(200)로 전송하는 단계(S130), 서버(200)가 글로벌 모델을 생성하는 단계(S150), 서버(200)가 글로벌 모델의 파라미터를 단말(100)로 전송하는 단계(S170) 및 단말이 글로벌 모델의 파라미터에 기초하여 로컬 모델을 업데이트하는 단계(S190)를 포함할 수 있다.
- [0119] 복수의 제2 로컬 모델을 생성하는 단계(S110)에서는, 단말(100)이 기 저장된 제1 로컬 모델로부터 복수의 제2 로컬 모델(LM)을 생성할 수 있다. 그리고 복수의 제2 로컬 모델을 생성하는 단계(S110)에서는, 제1 로컬 모델로부터 너비 변경이 가능한 인공지능 알고리즘을 기초로 복수의 제2 로컬 모델(LM)을 생성하되, 복수의 제2 로컬 모델(LM)의 너비는 제1 로컬 모델의 너비 이하일 수 있다.
- [0120] 이러한 복수의 제2 로컬 모델(LM)은, 제1 로컬 모델의 절반 너비를 갖는 하프(Half) 제2 로컬 모델 및 제1 로컬 모델의 전체 너비를 갖는 풀(full) 제2 로컬 모델 중 적어도 하나를 포함할 수 있다.
- [0121] 그리고 하프 제2 로컬 모델은, 제1 로컬 모델의 왼쪽 절반에 대응되는 LH-제2 로컬 모델(LH) 및 제1 로컬 모델의 오른쪽 절반에 대응되는 RH-제2 로컬 모델(RH)을 포함할 수 있다.
- [0122] 단말(100)이 복수의 제2 로컬 모델(LM)의 파라미터를 중첩 코딩하여 서버(200)로 전송하는 단계(S130)에서는 단말(100)이 통신 채널의 전송 전력을 일정 비율로 분배하고, 중첩 코딩된 복수의 제2 로컬 모델(LM)의 파라미터를 서로 다른 전송 전력으로 전송할 수 있다.
- [0123] 한편 서버(200)가 글로벌 모델을 생성하는 단계(S150)에서는, 서버(200)가 중첩 코딩된 복수의 제2 로컬 모델(SC-LM)의 파라미터를 디코딩할 수 있고, 디코딩된 파라미터를 집계(aggregates)하여 글로벌 모델(GM)을 생성할 수 있다.
- [0124] 그리고 이러한 글로벌 모델을 생성하는 단계(S150)에서는 서버(200)가 중첩 코딩된 복수의 제2 로컬 모델(SC-LM)의 파라미터를 연속적으로 디코딩할 수 있다. 이때 서버(200)에서의 복수의 제2 로컬 모델(LM)의 파라미터에 대한 획득여부는 통신 채널의 페이딩(fading) 이득에 기초하여 산출되는 디코딩 성공 확률에 따라 결정될 수 있다.
- [0125] 이에 글로벌 모델을 생성하는 단계(S150)에서는, 디코딩 성공 확률에 따라 LH-제2 로컬 모델(LH) 및 RH-제2 로컬 모델(RH) 중 하나의 파라미터만을 획득하거나 풀 제2 로컬 모델의 파라미터를 획득할 수 있다.
- [0126] 이후 서버(200)가 글로벌 모델의 파라미터를 단말(100)로 전송하는 단계(S170)가 수행될 수 있다.
- [0127] 그리고 단말(100)이 글로벌 모델의 파라미터에 기초하여 제1 로컬 모델을 업데이트하는 단계(S190)가 수행된다.
- [0128] 이러한 로컬 모델을 업데이트하는 단계(S190)에서 단말(100)은 복수의 제2 로컬 모델(LM) 중 하프 제2 로컬 모델과 같이 제1 로컬 모델보다 작은 너비를 갖는 제2 로컬 모델(LM)이 제1 로컬 모델과 유사한 로짓(logit)을 산출하도록 하는 인플레이스 지식 정제(IPKD, inplace knowledge distillation)를 제1 로컬 모델에 적용하여 제1 로컬 모델을 업데이트할 수 있다.
- [0130] 이하에서는 도 8 내지 도 10을 참조하여, 본 발명의 연합 학습 시스템의 효율성과 산업상 이용 가능성을 보여주는 실험의 수행과정 및 결과를 설명하기로 한다.
- [0131] 그리고 본 발명의 효과를 입증하기 위해 중첩 코딩 및 연속 디코딩을 사용하지 않는 연합 학습 시스템(이하, vanilla FL)을 본 발명의 연합 학습 시스템과 비교하였다.
- [0132] 실험을 위해 본 발명의 연합 학습 시스템은 중첩 코딩 및 연속 디코딩이 있는 다중 너비 인공신경망을 활용하여 장치 당 대역폭 W , 업링크 전송 전력 P 를 소비하여 0.5x 및 1.0x 모델을 동시에 교환하고 학습하였다. 이는 종래의 vanilla FL의 1.5x와 비교된다. vanilla FL은 너비 조정이 가능한 다중 너비 인공신경망이 없기 때문에 vanilla FL의 각 장치는 고정 너비 0.5x 및 1.0x 모델(각각 vanilla FL-0.5x 및 vanilla FL-1.0x)을 실행한다. 또한 중첩 코딩 또는 연속 디코딩이 없으므로 vanilla FL의 장치는 0.5x 및 1.0x 모델을 별도로 교환한다.
- [0133] 간단히 말해서 vanilla FL-1.5는 대역폭, 전송 전력 및 컴퓨팅 리소스를 두 배로 늘려 0.5x 및 1.0x 모델에 대

해 별도로 두 개의 연합 평균 연합을 동시에 실행하는 것과 같다.

[0134] 한편 시뮬레이션을 위해 Fashion MNIST 데이터 세트를 사용하여 기본적으로 분류 작업을 고려하였다. 그리고 장치 전반에 걸친 데이터 세트 분포의 non-IIDness는 농도 매개변수 $\alpha \in \{0.1, 1.0, 10\}$ 을 사용하는 Dirichlet 분포에 의해 제어되며, 여기서 α 가 낮을수록 IID가 더 많이 분포된다. 즉 장치 전반에 걸쳐 레이블에 비해 더 많은 불균형 샘플 수를 갖는 것으로 도 8과 같이 시각화 된 것이다. 도 8의 (a)는 non-IID로 $\alpha = 0.1$ 이고 도 8의 (b)는 IID로 $\alpha = 10$ 이다.

[0135] 그리고 업링크 및 다운링크 통신의 단일 라운드 이후에는 모든 단일 로컬 학습 에폭(epoch)이 수행된다. 서로 다른 장치를 통한 통신 채널은 업링크와 다운링크 모두에서 직교한다. 그리고 각 채널 구현에 대한 소규모 페이딩 이득 g 는 지수 분포 $g \sim \text{Exp}(1)$, 즉 레일리(Rayleigh) 페이딩을 따른다.

[0136] 통신 하이퍼파라미터는 하기의 표 1과 같다.

[0137] [표 1]

Description	Value
Initial learning rate (η_0)	10^{-3}
Optimizer	Adam
Distance (d)	100 [m]
Path loss exponent (β)	2.5
Bandwidth per device (W)	75×10^6 [Hz]
Central frequency (f_c)	5.9 [GHz]
Uplink transmission power (P)	23 [dBm]
Noise power spectrum (N_0)	-169 [dB/Hz]

[0139] 도 9의 (b)에 도시된 바와 같이 본 발명에 따른 연합 학습 시스템에서의 정확도는 연합 장치, 즉 단말의 수가 증가할수록 향상된다. 0.5x는 최대 79%의 정확도를 달성하였고, 1.0x는 최대 85%의 정확도를 달성하였다.

[0140] 또한 non-IID와 70개 이상의 단말을 가진 0.5x 연합 학습 시스템은 20개의 단말을 가진 1.0x 연합 학습 시스템보다 더 높은 정확도를 보여줄 수 있다.

[0141] 이러한 실험 결과를 바탕으로 non-IID 데이터 셋을 기반으로 연합 학습 시스템을 구성할 때 단말의 수와 다중 너비 인공지능망을 통한 너비를 조정함으로써 최적화를 달성할 수 있음을 알 수 있다.

[0142] 한편 다양한 통신 조건과 non-IID 데이터 분포에 있는 환경에서 종래의 vanilla FL과 비교하여 본 발명의 연합 학습 시스템의 성능을 검증하기 위한 실험을 수행하였다.

[0143] Non-IID 데이터에 대한 견고성(Robustness)은 도 10의 (d), (e), (f)와 하기의 표 2를 통해 알 수 있듯이 본 발명의 연합 학습 시스템인 SlimFL-0.5x는 α 의 조건에서 안정적인 수렴을 보인다.

[0144] [표 2]

Method	Top-1 Accuracy (%)					
	$\alpha = 0.1$	Good $\alpha = 1$	$\alpha = 10$	$\alpha = 0.1$	Poor $\alpha = 1$	$\alpha = 10$
SlimFL-0.5x	54 ± 2.2	83 ± 1.0	85 ± 1.0	56 ± 2.4	82 ± 1.7	85 ± 1.1
SlimFL-1.0x	59 ± 2.3	85 ± 1.1	87 ± 1.1	65 ± 2.9	84 ± 1.4	87 ± 0.9
Vanilla FL-0.5x	45 ± 5.9	84 ± 1.1	85 ± 1.0	39 ± 8.3	83 ± 1.2	85 ± 0.9
Vanilla FL-1.0x	69 ± 5.8	85 ± 4.0	86 ± 4.3	55 ± 9.2	80 ± 6.0	82 ± 4.7

[0146] 종래 기술인 vanilla FL-0.5x 및 vanilla FL-1.0x는 채널 상태가 불량하고, $\alpha=0.1$ 인 non-IID 데이터 분포 상태에서 8.3 및 9.2의 std를 나타낸다. 한편 본 발명의 연합 학습 시스템을 이용하는 slimFL-0.5x 및 slimFL-1.0x는 Top-1 Accuracy에서 2.4와 2.9의 std를 보인다. 이러한 경향은 $\alpha = 1$, $\alpha = 10$ 일 때도 유지된다.

[0147] 또한 본 발명의 연합 학습 시스템을 이용하는 slimFL-0.5x 및 slimFL-1.0x는 종래 기술인 vanilla FL-0.5x 및 vanilla FL-1.0x보다 낮은 변형을 보이는데, 이는 불량 채널의 non-IID 데이터에 대한 견고성을 강조한다.

[0148] 한편 도 10과 표 2는 본 발명의 연합 학습 시스템을 이용하는 slimFL 및 종래의 vanilla FL은 모두 양호한 채널 조건에서는 높은 정확도를 달성하는 것을 보여준다.

[0149] 하지만 도 10의 (c) 도 10의 (f)와 표 2와 같이 채널 상태가 양호에서 불량으로 악화됨에 따라 $\alpha = 10$ 일 때 vanilla FL-1.0x의 최대 정확도가 86%에서 82%로 떨어지는 것이 확인된다.

[0150] 한편 본 발명의 연합 학습 시스템을 이용하는 slimFL-1.0x의 정확도는 $\alpha = 10$ 일 때 양호 채널과 불량 채널 모두에서 동일한 최대 정확도 87%를 유지한다.

[0151] 또한 $\alpha = 0.1$ 일 때 slimFL-1.0x는 통신 및 컴퓨팅 비용을 더 많이 소요하는 vanilla FL-1.0x보다 18% 더 높은 Top-1 Accuracy를 달성한다.

[0152] 그리고 vanilla FL-1.0x의 Top-1 Accuracy의 규격은 채널 조건이 악화될수록 최대 59% 증가하는 반면, 본 발명의 연합 학습 시스템의 규격은 최대 31%만 증가함을 보였다. 이러한 결과는 열악한 채널에 대한 본 발명의 연합 학습 시스템의 견고성과 non-IID 데이터 배포(낮은 α) 및 통신 효율성에 대한 견고성을 보장한다.

[0153] 한편 이상적인 채널 조건(즉, 항상 성공적인 디코딩)에서 10개의 장치와 서버 간에 전송되는 총 비트량은 본 발명의 연합 학습 시스템인 slimFL과 vanilla FL-1.0x의 경우 205.8MBytes이고, vanilla FL-0.5x의 경우에는 102.9MBytes이다.

[0154] 하기 표 3은 본 발명의 slimFL이 중첩 코딩 및 연속 디코딩을 사용한 덕분에 vanilla FL-1.0x보다 최대 3.52% 적은 비트 손실을 달성함을 보여준다

[0155] [표 3]

Decoding Success Bits [MBytes]	0.5x	SlimFL 1.0x	drop	Vanilla FL-0.5x 0.5x	drop	Vanilla FL-1.0x 1.0x	drop
$\sigma^2 = -30\text{dB}$	1.96	198.45	5.46	102.21	0.72	200.30	5.56
$\sigma^2 = -40\text{dB}$	18.32	130.10	57.44	93.87	9.06	144.93	60.96

[0156]

[0157] slimFL의 감소된 드롭 비트는 vanilla FL-1.0x에서 동시에 수신할 수 없는 0.5x 모델의 성공적으로 디코딩된 비트에서 찾을 수 있다. slimFL의 전송 전력의 일부가 0.5x 로컬 모델에 할당되기 때문에 slimFL은 vanilla FL-1.0x보다 1.0x 모델 비트를 더 적게 디코딩한다.

[0158] 그 대가로 본 발명의 slimFL은 1.0x 로컬 모델뿐만 아니라 0.5x 로컬 모델도 동시에 받고, 추가로 받은 0.5x 로컬 모델은 1.0x 모델의 LH 세그먼트에 해당하므로 0.5x 및 1.0x 로컬 모델의 정확도와 수렴 속도가 향상된다.

[0159] 이에 본 발명의 slimFL은 이상의 이점을 누리면서도 vanilla FL-1.5x에 비해 절반의 전송 전력과 대역폭을 소비하므로 표 3에서 볼 수 있는 것처럼 통신 효율성을 높일 수 있다.

[0160] 이상에서는 고정된 1000 에폭으로 학습한 후 본 발명의 성능을 측정하였다. 이때 수렴까지 소비되는 에너지를 측정하는데, 학습 수렴은 테스트 정확도의 표준 편차(std)가 목표 임계 값 미만이고 최소 테스트 정확도가 100 개의 연속 라운드에서 평균 테스트 정확도보다 높아지는 순간으로 정의하였다.

[0161] 그리고 모델의 수렴을 측정하기 위해 평균 μ_{Ref} 의 참조값을 80%로, σ_{Ref} 를 7.2%로 정의하였다.

[0162] 100개의 연속 에폭에 대한 Top-1 Accuracy의 평균이 μ_{Ref} 보다 높고 평균 std가 σ_{Ref} 보다 낮을 때 수렴을 정의한다.

[0163] 그리고 표 4의 라운드 당 통신 및 컴퓨팅 에너지 비용을 감안할 때 수렴까지 본 발명의 연합 학습 시스템인 slimFL과 vanilla FL-1.5x의 총 에너지 비용을 비교한 결과는 표 5와 같다.

[0165] [표 4]

Metric	SlimFL	Vanilla FL-1.5x
Communication Cost [mW/Round]	199.5	399.1
Computation Cost [MFLOPS/Epoch]	3.56	3.56

[0166]

[0167] [표 5]

Metric	non-IIDness	SlimFL		Vanilla FL-1.5x	
		Good	Poor	Good	Poor
Communication Cost [W]	$\alpha = 0.1$	71.0	57.3	158.8	196.8
	$\alpha = 1.0$	8.5	10.4	15.8	36.7
	$\alpha = 10$	3.03	3.51	10.2	25.4
Computation Cost [GFLOPS]	$\alpha = 0.1$	1.27	1.02	1.88	2.41
	$\alpha = 1.0$	0.15	0.18	0.22	0.51
	$\alpha = 10$	0.05	0.06	0.14	0.35

[0168]

[0169] 결과는 평균적으로 본 발명의 연합 학습 시스템이 수렴될 때까지 총 통신 비용을 2.9배 낮추고 총 컴퓨팅 비용을 3.6배 절감함을 보여준다.

[0170] 이러한 더 높은 에너지 효율성의 결과는 본 발명의 연합 학습 시스템이 중첩 코딩 및 연속 디코딩에 의해 non-IID 및/또는 열악한 채널 조건에서도 더 빠른 수렴 속도를 보이기 때문에 가능하다.

[0171] 따라서 기존의 연합 학습 시스템에서는 사용가능한 에너지 및 채널 처리량이 서로 다른 다양한 기기종 장치에 대해 유연하게 대처할 수 없다는 문제를 본 발명의 연합 학습 시스템에서는 해결할 수 있게 된다.

[0172] 이는 본 발명의 연합 학습 시스템이 너비 변경이 가능한 다중-너비 인공신경망을 사용하는 장치에서의 학습을 위한 학습 알고리즘을 개발하고, 학습된 모델의 집계를 위해 중첩 코딩을 이용하여 서버로 전송하며, 서버에서 연속 디코딩을 통해 글로벌 모델을 생성할 수 있기 때문이다.

[0173] 특히나 광범위한 실험을 통해 본 발명의 연합 학습 시스템은 다양한 통신 환경과 데이터 배포 환경에서 통신 및 에너지를 효율적으로 활용할 수 있는 솔루션임을 확인하였다.

[0174] 그리고 이와 같은 본 발명의 연합 학습 방법은 다양한 컴퓨터 구성요소를 통하여 수행될 수 있는 프로그램 명령어의 형태로 구현되어 컴퓨터 판독 가능한 기록 매체에 기록될 수 있다. 상기 컴퓨터 판독 가능한 기록 매체는 프로그램 명령어, 데이터 파일, 데이터 구조 등을 단독으로 또는 조합하여 포함할 수 있다.

[0175] 상기 컴퓨터 판독 가능한 기록 매체에 기록되는 프로그램 명령어는 본 발명을 위하여 특별히 설계되고 구성된 것들이거나 컴퓨터 소프트웨어 분야의 당업자에게 공지되어 사용 가능한 것일 수도 있다.

[0176] 컴퓨터 판독 가능한 기록 매체의 예에는, 하드 디스크, 플로피 디스크 및 자기 테이프와 같은 자기 매체, CD-ROM, DVD 와 같은 광기록 매체, 플롭티컬 디스크(floptical disk)와 같은 자기-광 매체(magneto-optical media), 및 ROM, RAM, 플래시 메모리 등과 같은 프로그램 명령어를 저장하고 수행하도록 특별히 구성된 하드웨어 장치가 포함된다.

[0177] 프로그램 명령어의 예에는, 컴파일러에 의해 만들어지는 것과 같은 기계어 코드뿐만 아니라 인터프리터 등을 사용해서 컴퓨터에 의해서 실행될 수 있는 고급 언어 코드도 포함된다. 상기 하드웨어 장치는 본 발명에 따른 처리를 수행하기 위해 하나 이상의 소프트웨어 모듈로서 작동하도록 구성될 수 있으며, 그 역도 마찬가지이다.

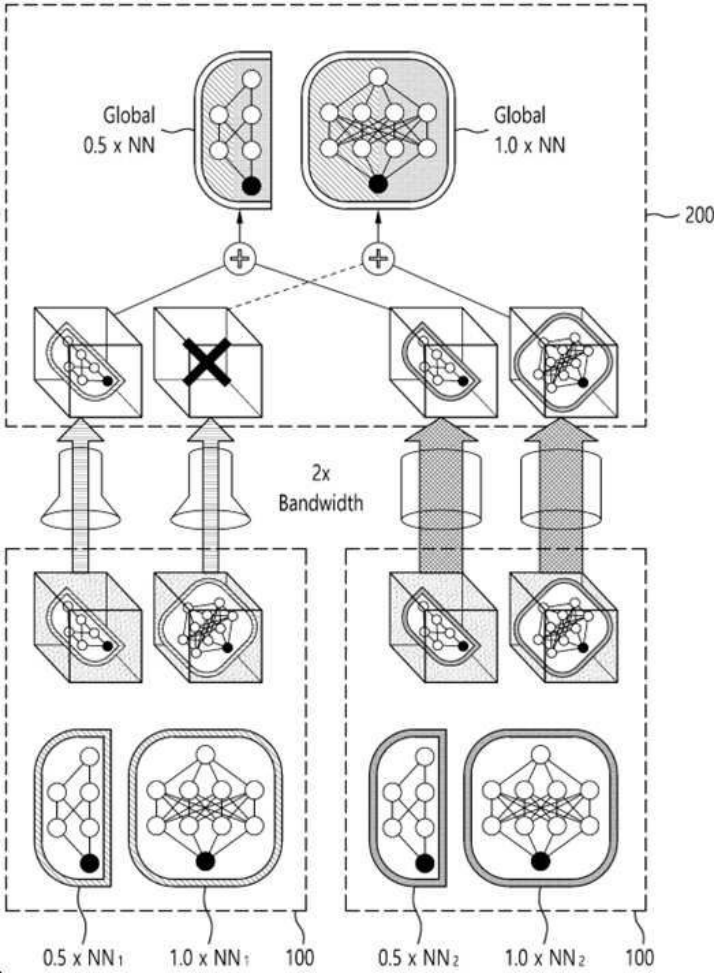
[0178] 이상에서는 본 발명의 다양한 실시예에 대하여 도시하고 설명하였지만, 본 발명은 상술한 특징의 실시예에 한정되지 아니하며, 청구범위에서 청구하는 본 발명의 요지를 벗어남이 없이 당해 발명이 속하는 기술분야에서 통상의 지식을 가진 자에 의해 다양한 변형실시가 가능한 것은 물론이고, 이러한 변형실시들은 본 발명의 기술적 사상이나 전망으로부터 개별적으로 이해되어져서는 안 될 것이다.

부호의 설명

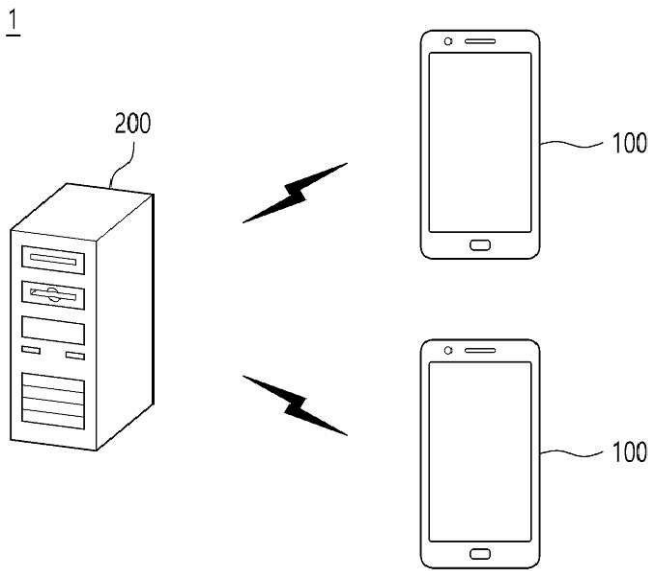
[0179] 1 : 연합 학습 시스템 100 : 단말
200 : 서버

도면

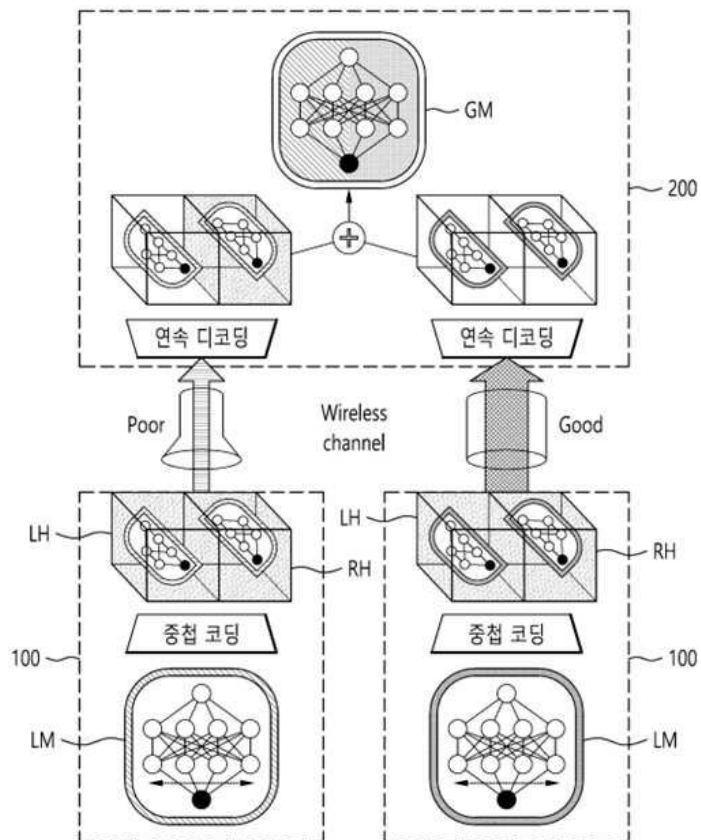
도면1



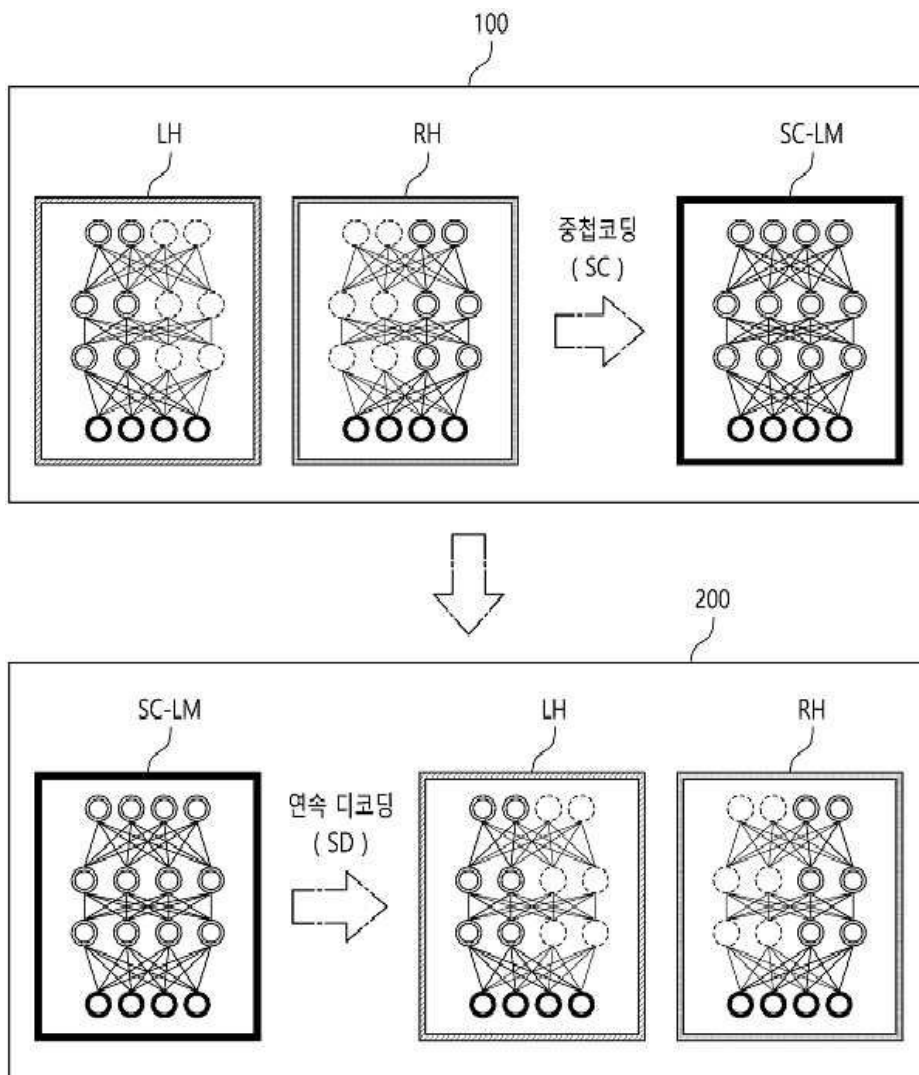
도면2



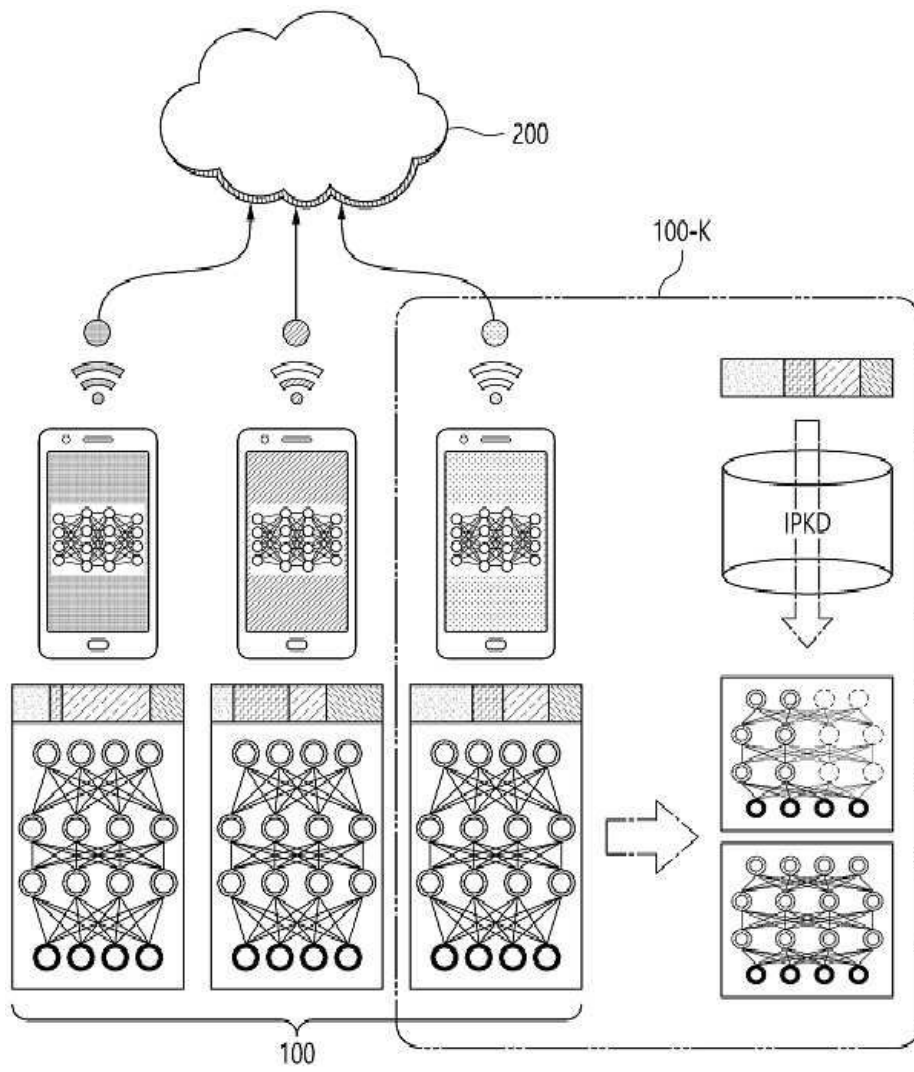
도면3



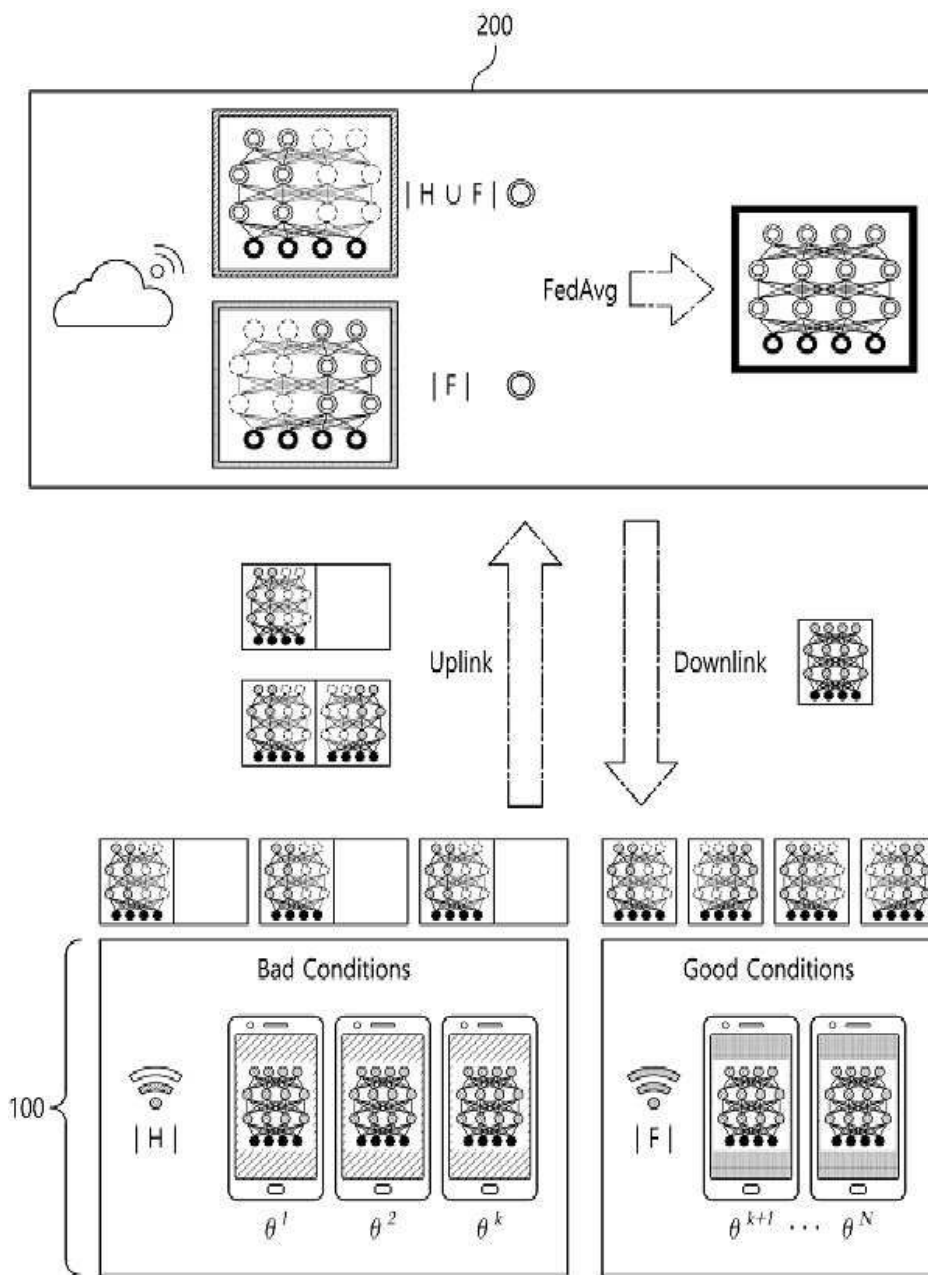
도면4



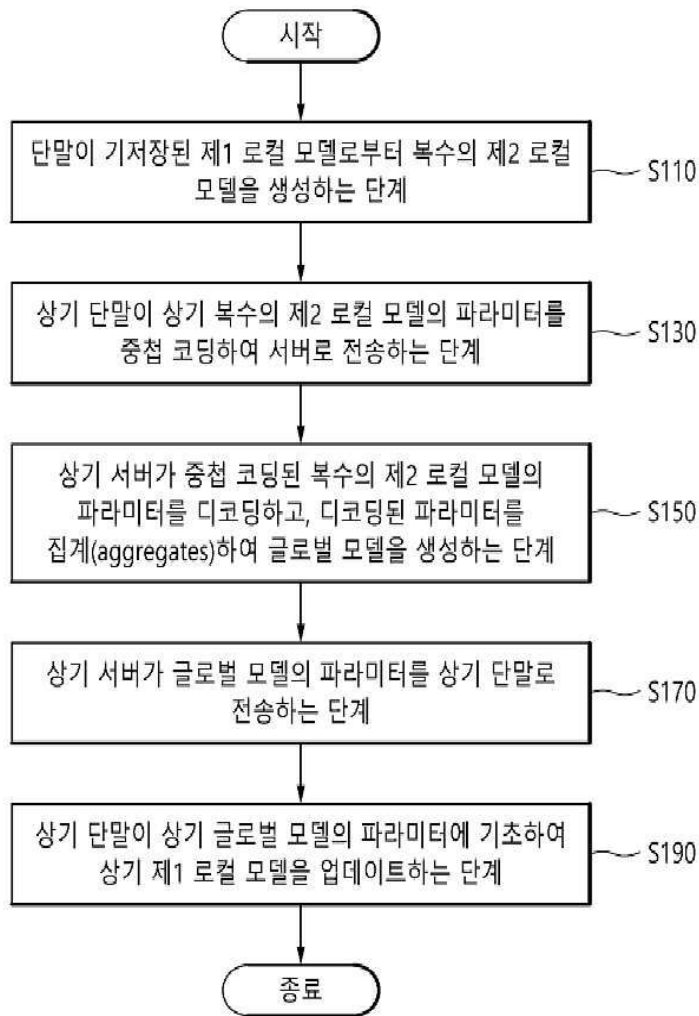
도면5



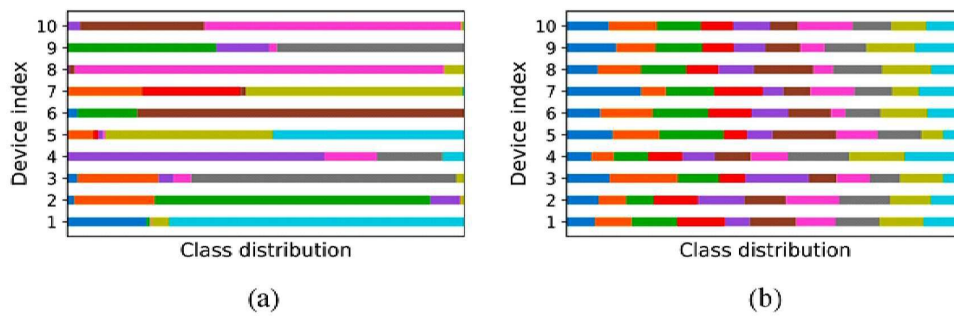
도면6



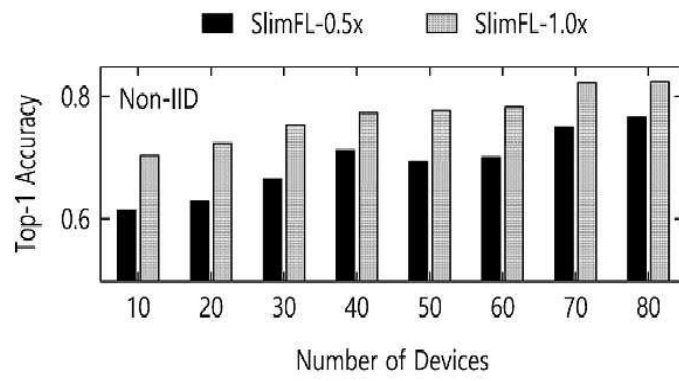
도면7



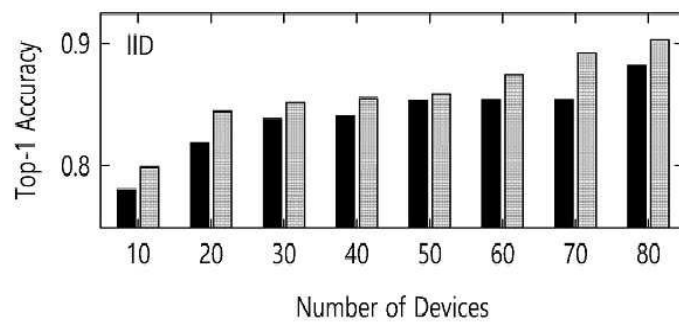
도면8



도면9



(a)



(b)

도면10

