

面向网络贸易交易的联邦学习最优委托策略分析

张雨豪, 曾 进, 杜前程

贵州大学数学与统计学院, 贵州 贵阳

收稿日期: 2024年3月1日; 录用日期: 2024年3月14日; 发布日期: 2024年5月23日

摘 要

近年来, 数字化经济极大地推动了网络贸易交易的发展。联邦学习技术作为一种分布式的机器学习范式, 能够助力网络中分布式的节点进行贸易合作。然而, 分布式节点与云服务器之间具备信息不对称性, 现有关于联邦学习场景的相关方案存在交互模型构建不准确问题, 导致云服务器难以选择最优委托策略。所提出的方法基于信号博弈模型对联邦学习进行建模, 并对该场景下的参与者交互问题进行分析。首先, 对参与者进行理性假设, 定义参与者集合、行动空间、支付函数等博弈要素。其次, 引入行为 - 信念系统, 构建服务器与用户之间的信号博弈模型。最后, 由序贯理性和序贯一致性, 证明该博弈存在序贯均衡, 求解服务器针对不同类型用户的最优委托策略。通过实验仿真, 验证选择最优委托策略的必要性。

关键词

联邦学习, 信息不对称, 网络贸易交易, 最优委托策略

Analysis of Optimal Delegation Strategies for Federated Learning in the Context of Online Trade Transactions

Yuhao Zhang, Jin Zeng, Qiancheng Du

School of Mathematics and Statistics, Guizhou University, Guiyang Guizhou

Received: Mar. 1st, 2024; accepted: Mar. 14th, 2024; published: May 23rd, 2024

Abstract

In recent years, the digital economy has greatly promoted the development of online trade trans-

文章引用: 张雨豪, 曾进, 杜前程. 面向网络贸易交易的联邦学习最优委托策略分析[J]. 电子商务评论, 2024, 13(2): 2130-2141. DOI: 10.12677/ecl.2024.132259

actions. As a distributed machine learning paradigm, federated learning technology can help distributed nodes in the network to conduct trade cooperation. However, there is information asymmetry between distributed nodes and cloud servers. The existing related schemes about federated learning scenarios have the problem of inaccurate interaction model construction, which makes it difficult for cloud servers to choose the optimal delegation strategy. The proposed method models federated learning based on the signal game model, and analyzes the problem of participant interaction in this scenario. Firstly, the rational hypothesis of the participants is carried out, and the game elements such as the set of participants, the action space and the payment function are defined. Secondly, the behavior-belief system is introduced to construct the signal game model between the server and the user. Finally, by sequential rationality and sequential consistency, it is proved that the game has sequential equilibrium, and the optimal delegation strategy of the server for different types of users is solved. Through experimental simulation, the necessity of selecting the optimal delegation strategy is verified.

Keywords

Federated Learning, Information Asymmetry, Network Trade Transactions, Optimal Delegation Strategy

Copyright © 2024 by author(s) and Hans Publishers Inc.

This work is licensed under the Creative Commons Attribution International License (CC BY 4.0).

<http://creativecommons.org/licenses/by/4.0/>



Open Access

1. 引言

近年来,随着数字化经济的快速发展,网络贸易已成为全球经济中不可或缺的一部分。在该趋势下,机器学习技术日益被视为网络贸易领域的重要驱动力。然而,与传统行业相比,网络贸易交易的特点是涉及多方参与者、涵盖多种产品和服务,并伴随着大量复杂的数据交换和处理。然而,网络贸易中的各参与方的数据通常涉及交易细节、客户隐私等重要信息,往往具有高度敏感性,而传统的集中式数据处理方式存在安全隐患,可能导致数据泄露、信息篡改等风险[1]。因此,如何在保障数据安全的前提下实现数据共享和合作成为网络贸易交易中的关键问题之一[2]。

联邦学习[3]是一种新兴的分布式机器学习框架,为解决网络交易贸易中的数据安全和隐私保护问题提供了新的思路和方法。其允许每个设备或数据中心在本地训练模型,并将更新的模型参数上传至服务器。分布式训练的核心结构不仅使得联邦学习最大限度地保护了用户的数据隐私,避免了敏感信息泄露的风险,还减少了信道中的数据传输量,节省了服务器的计算资源。此外,分布式数据源的优势使得联邦学习训练出的模型性能和泛化能力提升[4] [5] [6]。

联邦学习的种种优势使其广泛应用于诸多领域。例如,在电子商务领域,各电商平台可以利用联邦学习共同训练推荐系统模型,提高用户购物体验 and 交易效率,同时保护用户隐私。在供应链管理中,各供应商可以利用联邦学习共同优化供应链预测模型,实现更高效的库存管理和物流运输。在医疗领域中,诊疗数据通常非常敏感,联邦学习可以帮助存在信息壁垒的医疗机构共同训练模型,进一步改善疾病诊断和治疗效率[7]。在金融领域中,银行、信用卡公司等金融机构通常需要处理大量的客户数据,而联邦学习能够帮助其提高客户推荐等方面的准确率,提高反欺诈率的同时保护客户隐私[8]。在交通系统中,联邦学习可以帮助多个城市可以共同训练交通拥堵预测模型,以提高路况预测的准确率,实现更加智能化的交通系统[9]。此外,物联网设备通常具有大量的传感器数据,联邦学习可以帮助将这些数据进行有

效的分析和利用,例如共同训练智能家居控制模型、智能健康监测等模型[10]。可见,在不同的组织和企业拥有私密敏感的数据集时,联邦学习可以帮助它们利用这些无法直接共享的数据,更好地释放数据潜在价值。

联邦学习充分发挥其潜力的先决条件在于能够激励高质量用户以真实方式参与合作训练[11],然而,每位用户的计算、通信和隐私成本均属于服务器无法获取的私有信息,由此产生了信息不对称,服务器无法预先了解各用户所拥有的计算资源。而当大量计算能力较低的用户参与到训练中,则可能引发模型收敛速度减缓、联邦学习效率下降以及高质量用户缺乏参与动力等难题[12][13][14],导致联邦学习中的协作难题。此外,现有的解决方案仍然存在以下问题:1) 没有考虑以算力的高低对用户进行区分。2) 缺乏对联邦学习信息不对称场景下的准确建模和问题形式的合理分析。

针对上述挑战,本文基于信号博弈构建一个合理、有效的交互博弈模型,以指导云服务器针对不同类型的用户选择最优委托策略,解决网络贸易交易中的数据安全和隐私保护问题。本文的主要工作为:

1) 针对联邦学习的信息不对称现象,构建一个适用于该信息不对称场景下的信号博弈模型,通过实际地考虑用户的计算资源差异,将参与者划分为高算力和低算力两种类型,解决现有研究在交互模型构建方面不准确的问题,以进一步帮助各方参与者有效地调整合作策略,提高系统效率。

2) 针对联邦学习中服务器难以确定最优委托策略的问题,求解信号博弈模型的序贯均衡,通过数值仿真验证选择最优委托策略对提高联邦学习效率的关键作用,有效地解决由于信息不对称所带来的服务器无法预先获取各用户计算资源信息的挑战以及由于信息不对称而导致的模型性能下降问题,为网络贸易交领域中制定各方合作策略提供有力支持。

2. 相关概念

2.1. 信号博弈模型

定义 1 信号博弈

1) 参与人集合: $\{0\} \cup N = \{1, 2, \dots, n\}$: 其中, $\{0\}$ 代表虚拟参与人“自然”。

2) 博弈树 $(X, <)$: 其中 $<$ 是一个在节点集合 X 上严格偏序贯性(非自反性, 传递性)的“先于”行动顺序。

3) 节点: 对于节点集合 X 中的元素定义其中一种节点的类型: 终结节点 Z , 即博弈结束的节点。

4) 参与人函数 I : 定义一个函数 $I: X \setminus Z \rightarrow N \cup \{0\}$, 即对于所有非终结节点(决策点), 函数 I 使决策点对应到每一个参与人。

5) 信息集 H_i : 对所有参与人 $\forall i \in N \cup \{0\}$, 有参与人 i 的所有决策点集合 $X_i = \{x \in X \mid I(x) = i\}$, 存在 $H_i \subseteq 2^{X_i}$ 对决策点集合 X_i 分割为互不相交的集合, H_i 成为信息集; 所有参与人的信息集 $H = \bigcup_{i \in N \cup \{0\}} H_i$ 。

6) 行动空间: 在任何的信息集 $\forall h \in H$ 上, 都存在一个行动空间 A_h :

$\forall x \in h, A_x = A_h$, 即在同一个信息集中, 不同的节点的行动空间是一样的;

参与人 i 所有行动的集合 $A_i = \bigcup_{h \in H_i} A_h$, 所有参与人的行动集合 $A = \bigcup_{i \in N \cup \{0\}} A_i = \bigcup_{h \in H} A_h$ 。

7) 先验概率分布

对于“自然”参与人, 它的信息集对应一个概率分布 $P: H_0 \rightarrow \bigcup_{h \in H_0} \Delta A_h$, 即信息集对应的行动空间上, 每个行动都有一个 0 到 1 之间的概率, 且所有的概率之和为 1。

8) 效用函数

对于参与人 $\forall i \in N$, 效用由博弈树中的终结节点 Z 确定, 即 $u_i: Z \rightarrow \mathbb{R}$ 。

2.2. 序贯均衡

定义 2 (行为—信念系统) 假设 $B = \times_{i \in N} B_i = \times_{i \in N} \times_{h \in H_i} \Delta(A_h)$, $\Pi = \times_{i \in N} \times_{h \in H_i} \Delta(h)$, 其中, B 中的所有元素 $\beta \in B$ 称为行为向量, 即对于所有参与者的行动空间上都存在一个概率分布。 Π 中的所有元素 $\pi \in \Pi$ 称为信念向量, 即对于所有参与者的每一个信息集上都存在一个概率分布。组合 $(\beta, \pi) \in B \times \Pi$ 称为行为—信念系统。

定义 3 (序贯均衡) 博弈中, $(\beta, \pi) \in B \times \Pi$ 是一个序贯均衡, 其中, 当且仅当其满足以下条件:

- 1) (β, π) 是序贯理性的: 如果任意参与者 $\forall i \in N$, 在其任意信息集 $\forall h \in H_i$ 上, 对于 $\forall m_h \in \Delta(A_h)$, 都有 $u_i(\beta_h, \beta_{-h} | h) \geq u_i(m_h, \beta_{-h} | h)$, 其中 $m_h \neq \beta_h$, 则称 (β, π) 是序贯理性的。
- 2) (β, π) 是序贯一致的: 存在一个序列 $(\beta^k, \pi^k)_{k=1,2,\dots}$, 使得对任意 $\forall k \in \{1, 2, \dots\}$, (β^k, π^k) 是满足贝叶斯法则的, 即当 $k \rightarrow \infty$ 时, (β^k, π^k) 是收敛于 (β, π) 的, 则称 (β, π) 是序贯一致的。

3. 系统模型

联邦学习系统由两个部分组成, 分别是一个云服务器和一组用户, 如图 1 所示。

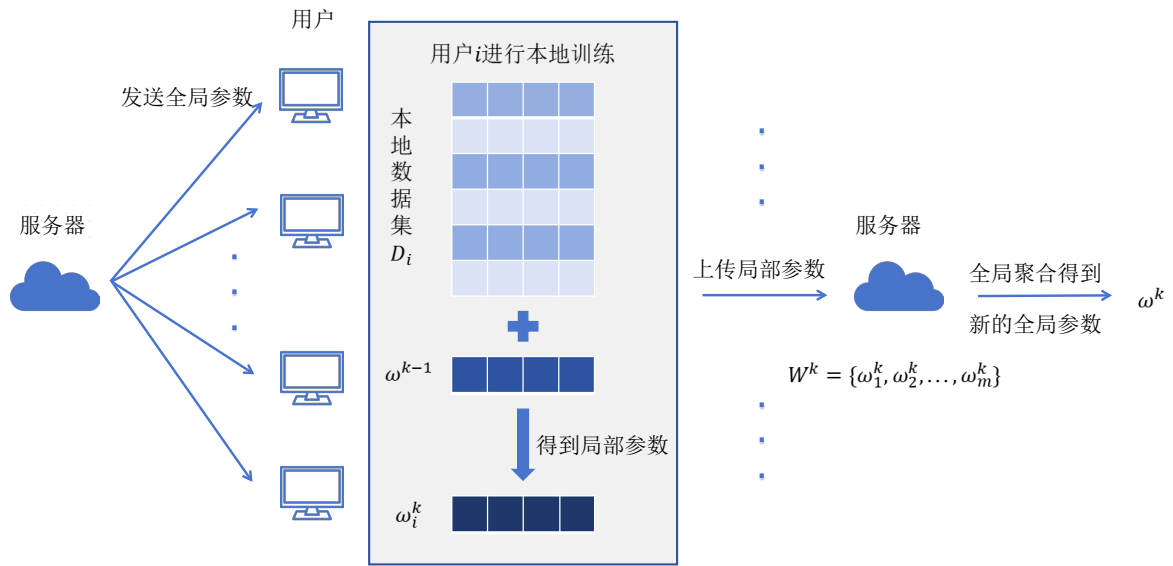


Figure 1. Federated learning system model

图 1. 联邦学习系统模型

用户集合: $\mathcal{M} = \{1, 2, \dots, m\}$, 服务器初始化模型参数 ω^0 后发送至每位用户。

用户 i 接收到参数 ω^{k-1} 后使用本地数据 D_i 通过梯度下降法更新局部模型参数 ω_i^k , 随后打包上传至服务器, 其中局部模型更新规则如式(1)所示。

$$\omega_i^k = \omega^{k-1} - \lambda \nabla L(\omega^{k-1}) \quad (1)$$

其中, $\lambda > 0$ 为学习率, $\nabla L(\omega^{k-1})$ 为损失函数的梯度。损失函数定义为:

$$L(\omega^{k-1}) = \frac{1}{|D_i|} \sum_{d_j \in D_i} l_j(\omega^{k-1}, d_j) \quad (2)$$

其中, $l_j(\omega^{k-1}, d_j)$ 是每个数据样本 $d_j \in D_i$ 的损失函数。

服务器接收到多方参数集合 $\mathcal{W}^k = \{\omega_1^k, \omega_2^k, \dots, \omega_n^k\}$ 后通过全局聚合得到新的全局模型 ω^k , 全局模型聚

合规则如式(3)所示。

$$\omega^k = \frac{1}{\sum_{i \in N} |D_i|} \sum_{i \in N} \omega_i^k \quad (3)$$

上述步骤将重复，直至模型参数收敛。本文涉及的符号和符号意义，如表 1 所示：

Table 1. Notations definition

表 1. 符号定义

符号	描述
D_i	用户 i 的本地数据集
k	全局训练的轮次
ω^k	第 k 轮训练的全局模型参数
ω_i^k	用户 i 在第 k 轮的本地模型参数
ω^0	初始全局模型参数
V	模型价值
W	服务器给予用户的报酬
C_β^α	训练成本， $\alpha = h, l$ ，分别表示高算力 and 低算力用户。 $\beta = H, L$ 分别表示高质量训练和低质量训练

4. 联邦学习信号博弈模型

结合联邦学习框架和定义 1，本文构造联邦学习场景下的信号博弈模型 $G = \langle N, T, \Theta, A, H, B, \Pi, U \rangle$ 。

- 1) 参与人集合为 $N = \{1, 2\}$ 。 $\{1\}$ 代表用户，作为信号的发送者， $\{2\}$ 代表服务器，作为信号的接收者。
- 2) 信号发送者的类型空间为 $T = \{\text{高算力用户}, \text{低算力用户}\}$ 。
- 3) 信号发送者的先验概率分布为 $\Theta = (\theta, 1 - \theta)$ 。该先验概率分布表示在“自然”情况下，服务器对用户类型的初始判断。
- 4) 参与人的行动空间为 $A = (A_1, A_2)$ 。 $A_1 = \{\text{高质量训练}, \text{低质量训练}\}$ ，表示用户的行动空间， $A_2 = \{\text{委托}, \text{不委托}\}$ 表示服务器的行动空间。
- 5) 博弈的信息集为 $H = (H_0, H_1, H_2, H_3, H_4)$ 。
- 6) 行动空间 A 上的对应的概率分布为 $B = (p, 1 - p; q, 1 - q; s, 1 - s; t, 1 - t)$ 。
- 7) 信息集 H 上对应的概率分布为 $\Pi = (1; 1; 1; \mu, 1 - \mu; \delta, 1 - \delta)$ ，也称为后验概率分布。服务器根据用户的行动，使用贝叶斯法则对先验概率 Θ 进行信念修正得到后验概率。
- 8) 参与人的效用函数为 $U = (U_1, U_2)$ 。其中， $U_1 = (u_{11}, u_{12}, \dots, u_{18})$ ， $U_2 = (u_{21}, u_{22}, \dots, u_{28})$ ，效用函数由所有参与者所选择的策略共同决定，参与者不同的策略组合的效用函数不同。

根据上述形式化定义，所构建的联邦学习信号博弈模型的博弈树形式，如图 2 所示。

5. 最优委托策略选取

本节由定义中序贯理性和序贯一致性，证明所构建的联邦学习信号博弈模型存在序贯均衡，求解服务器针对不同类型用户的最优委托策略。

在求解最优委托策略之前，本文规定：

- ◆ 无论是哪种类型的用户，高质量训练的成本恒大于低质量训练的成本，即 $C_H^h > C_L^h$ ， $C_H^l > C_L^l$ 。
- ◆ 对于参与人 1，终结节点 Z_5 与 Z_8 上的效用 $W - C_H^h < -C_L^h$ ，代表低算力用户需要支付价格补偿以换取更高质量的模型参数。其它终结节点上 $W - C_\beta^\alpha > 0$ 。

- ◆ 在实际的委托场景中，服务器与用户所采取的行动是确定的，不存在不确定性，故本文不讨论行动空间上概率不为 0 或 1 的情况。

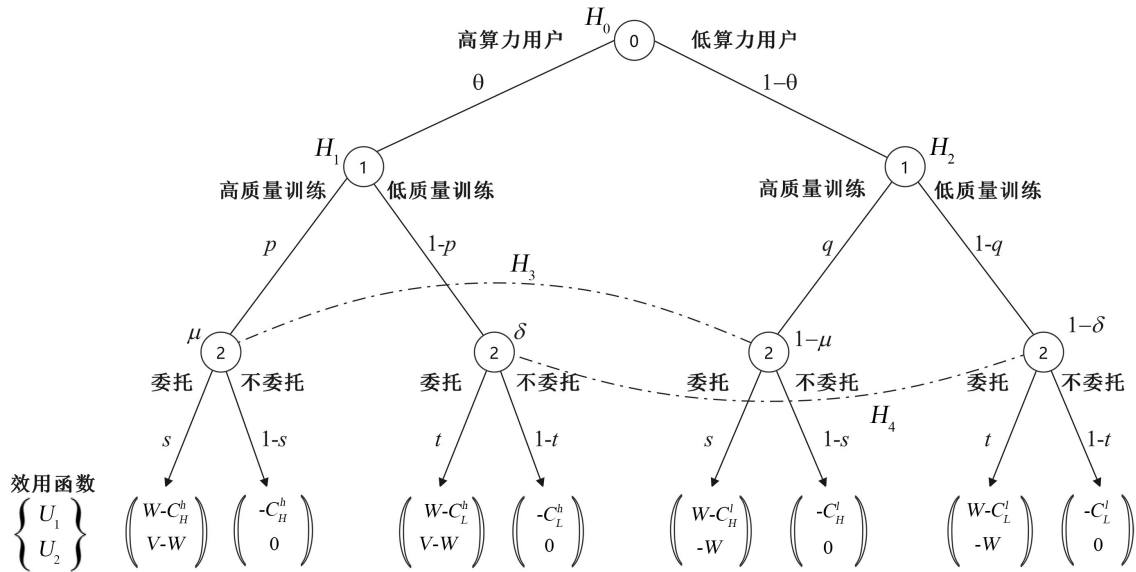


Figure 2. Signal game tree of federated learning

图 2. 联邦学习的信号博弈树

定理 1: 在博弈 G 中, 存在以下 4 种序贯均衡, 分别为:

均衡 1: $\{p=1, q=0, s=1, t=0, \mu=1, \delta=0\}$

均衡 2: 当 $\theta > \frac{W}{V}$ 时, $\left\{p=0, q=0, s=1, t=1, \mu \in \left[\frac{W}{V}, 1\right], \delta=\theta\right\}$

均衡 3: 当 $\theta > \frac{W}{V}$ 时, $\left\{p=0, q=0, s=0, t=1, \mu \in \left[0, \frac{W}{V}\right], \delta=\theta\right\}$

均衡 4: 当 $\theta < \frac{W}{V}$ 时, $\left\{p=0, q=0, s=0, t=0, \mu \in \left[0, \frac{W}{V}\right], \delta=\theta\right\}$

证明: 对于低算力用户来说, 不管服务器选择哪种策略, 低算力用户选择{低质量训练}的收益恒大于{高质量训练}的收益, 故在信息集 H_2 上, $q=0$ 恒成立。

由贝叶斯法则可得到式(4)和式(5)如下。

$$\mu = \frac{\theta p}{\theta p + (1-\theta)q} \quad (4)$$

$$\delta = \frac{\theta(1-p)}{\theta(1-p) + (1-\theta)(1-q)} \quad (5)$$

1) 在信息集 H_3 上, 假设服务器选择{委托}: $s=1$, 由序贯理性, 有

$$(V-W)\mu + (-W)(1-\mu) > 0 \cdot \mu + 0 \cdot (1-\mu), \text{ 即需满足 } \mu > \frac{W}{V}.$$

在信息集 H_4 上分以下两种情况讨论:

a) 假设服务器选择{不委托}: $t=0$, 由序贯理性, 有 $(V-W)\delta + (-W)(1-\delta) < 0 \cdot \delta + 0 \cdot (1-\delta)$, 即需

满足 $\delta < \frac{W}{V}$ 。

由逆向归纳法, 在信息集 H_1 上, 高算力用户选择{高质量训练}: $p=1$ 。

由公式(4)(5), 求得 $\mu=1$, $\delta=0$, 均满足序贯理性。

综上所述, $\{p=1, q=0, s=1, t=0, \mu=1, \delta=0\}$ 是一个序贯均衡, 记为均衡 1。

b) 假设服务器选择{委托}: $t=1$, 由序贯理性, 有 $(V-W)\delta + (-W)(1-\delta) > 0 \cdot \delta + 0 \cdot (1-\delta)$, 即需满足 $\delta > \frac{W}{V}$ 。

由逆向归纳法, 高算力用户选择{低质量训练}: $p=0$ 。

由公式(4), μ 为无穷小量, 证明其满足序贯一致性: 则存在一个趋于 0 的序列 $\tilde{\varepsilon}:(\varepsilon_1, \varepsilon_2)$, 当 $\tilde{\varepsilon} \rightarrow 0$ 时, $\mu^\varepsilon \rightarrow \mu$ 。

如果 $\mu \in \left(\frac{W}{V}, 1\right)$, 设 $\varepsilon_2 = c\varepsilon_1$, $c > 0$, 当 $\tilde{\varepsilon} \rightarrow 0$ 有:

$$\mu^\varepsilon = \frac{\theta\varepsilon_1}{\theta\varepsilon_1 + (1-\theta)c\varepsilon_1} = \frac{\theta}{\theta + (1-\theta)c} \quad (6)$$

即存在一个序列 $\tilde{\varepsilon}:(\varepsilon_1, \varepsilon_2)$, 当 $\varepsilon_2 = c\varepsilon_1$ 时 ($c = \frac{\theta}{1-\theta} \left(\frac{1}{\mu} - 1\right)$), 使得 $\mu \in \left(\frac{W}{V}, 1\right)$ 满足序贯一致。

如果 $\mu=1$, 设 $\varepsilon_2 = c\varepsilon_1^2$, $c > 0$, 当 $\tilde{\varepsilon} \rightarrow 0$ 时, 有:

$$\mu^\varepsilon = \frac{\theta\varepsilon_1}{\theta\varepsilon_1 + (1-\theta)c\varepsilon_1^2} = \frac{\theta}{\theta + (1-\theta)c\varepsilon_1} \quad (7)$$

即存在一个序列 $\tilde{\varepsilon}:(\varepsilon_1, \varepsilon_2)$, 当 $\varepsilon_2 = c\varepsilon_1^2$ 时 (c 取任意大于 0 的值), 使得 $\mu=1$ 满足序贯一致。

由公式(5), 可得 $\delta = \theta$, 故当 $\theta > \frac{W}{V}$ 时, 以上假设才可成立。

综上所述, $\left\{p=0, q=0, s=1, t=1, \mu \in \left(\frac{W}{V}, 1\right], \delta = \theta\right\}$ 是一个序贯均衡, 记为均衡 2。

2) 在信息集 H_3 上, 假设服务器选择{不委托}: $s=0$, 由序贯理性, 有

$(V-W)\mu + (-W)(1-\mu) < 0 \cdot \mu + 0 \cdot (1-\mu)$, 即需满足 $\mu < \frac{W}{V}$ 。

由于 $q=0$, 由公式(4)可知, $\mu = \frac{\theta p}{\theta p + (1-\theta)q}$ 为无穷小量, 若使 $\mu < \frac{W}{V}$, 则 p 的值必须取 0。

证明其满足序贯一致性: 则存在一个趋于 0 的序列 $\tilde{\varepsilon}:(\varepsilon_1, \varepsilon_2)$, 当 $\tilde{\varepsilon} \rightarrow 0$ 时, 有 $\mu^\varepsilon \rightarrow \mu$ 。

如果 $\mu \in \left(0, \frac{W}{V}\right)$, 设 $\varepsilon_2 = c\varepsilon_1$, $c > 0$, 当 $\tilde{\varepsilon} \rightarrow 0$ 有:

$$\mu^\varepsilon = \frac{\theta\varepsilon_1}{\theta\varepsilon_1 + (1-\theta)c\varepsilon_1} = \frac{\theta}{\theta + (1-\theta)c} \quad (8)$$

即存在一个序 $\tilde{\varepsilon}:(\varepsilon_1, \varepsilon_2)$, 当 $\varepsilon_2 = c\varepsilon_1$ 时 ($c = \frac{\theta}{1-\theta} \left(\frac{1}{\mu} - 1\right)$), 使得 $\mu \in \left(0, \frac{W}{V}\right)$ 满足序贯一致。

如果 $\mu=0$, 设 $\varepsilon_2 = c\sqrt{\varepsilon_1}$, $c > 0$, 当 $\tilde{\varepsilon} \rightarrow 0$ 时, 有:

$$\mu^\varepsilon = \frac{\theta\varepsilon_1}{\theta\varepsilon_1 + (1-\theta)c\sqrt{\varepsilon_1}} = \frac{\theta\sqrt{\varepsilon_1}}{\theta\sqrt{\varepsilon_1} + (1-\theta)c} \quad (9)$$

即存在一个序列 $\tilde{\varepsilon}:(\varepsilon_1, \varepsilon_2)$, 当 $\varepsilon_2 = c\sqrt{\varepsilon_1}$ 时 (c 取任意大于 0 的值), 使得 $\mu=0$ 满足序贯一致。

在信息集 H_4 上, 由公式(5)可得, $\delta = \theta$ 。

服务器选择{委托}的期望效用为:

$$(V - W)\delta + (-W)(1 - \delta) = \theta V - W \tag{10}$$

服务器选{不委托}的期望效用为:

$$0 \cdot \delta + 0 \cdot (1 - \delta) = 0 \tag{11}$$

- 1) 当 $\theta V - W > 0$, 即 $\theta > \frac{W}{V}$ 时, 服务器选择{委托}的期望效用大于选择{不委托}的期望效用, 此时 $t = 1$, 综上所述, $\left\{ p = 0, q = 0, s = 0, t = 1, \mu \in \left[0, \frac{W}{V} \right), \delta = \theta \right\}$ 是一个序贯均衡, 记为均衡 3。
- 2) 当 $\theta V - W < 0$, 即 $\theta < \frac{W}{V}$ 时, 服务器选择{委托}的期望效用小于选择{不委托}的期望效用, 此时 $t = 0$, 综上所述, $\left\{ p = 0, q = 0, s = 0, t = 0, \mu \in \left[0, \frac{W}{V} \right), \delta = \theta \right\}$ 是一个序贯均衡, 记为均衡 4。
- 证毕。

6. 仿真分析

该章节中, 对本文的博弈模型进行了仿真, 结果验证了博弈均衡解的存在和对应存在条件, 在每一种模拟情况下固定了一些参数, 相关收益和成本的取值皆根据第 4 章所规定的大小关系进行设定, 默认参数如表 2 所示。

Table 2. Experimental parameter settings
表 2. 实验参数设置

参数	大小设置
W	10
V	20
θ	0.5
C_H^h	5
C_L^h	3
C_H^L	11
C_L^L	5

首先, 固定其它参数的默认值, 只改变先验概率 θ , 观察了在四种不同的均衡策略下, 先验概率 θ 对服务器效用的影响, 结果如图 3 所示。

当 $\theta < \frac{W}{V} = 0.5$ 时, 存在均衡 1 与均衡 4, 且此时先验概率 θ 的变化并不会使服务器的收益大小发生改变。当 $\theta > \frac{W}{V} = 0.5$ 时, 存在均衡 1、均衡 2 及均衡 3, 对于均衡 2 和均衡 3 来说, 服务器的收益随着高算力用户比例的升高而增加, 而均衡 1 中服务器的收益始终不变。

由此可见, 不论先验概率 θ (高算力用户的比例)如何变化, 均衡 1 始终存在, 并且在此均衡策略下服务器收益始终大于其他均衡下的收益, 说明在任何“自然”情况下, 服务器在接收到信号为{高质量训练的模型参数}时, 服务器认为该用户为高算力用户($\mu = 1$), 并且会选择{委托}该用户, 并给予报酬; 服务器在接收到信号为{低质量训练}的模型参数时, 服务器认为该用户为低算力用户($\delta = 0$), 会选择{不委托}该用户。

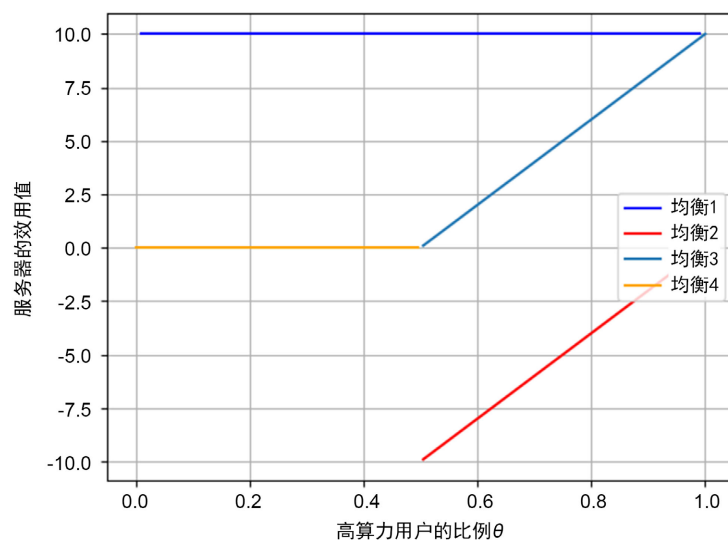
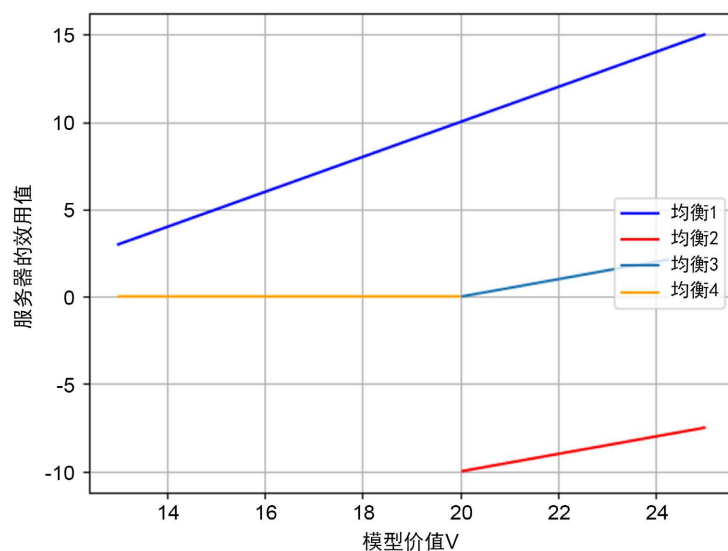


Figure 3. The influence of prior probability on server utility under four different equilibrium strategies

图 3. 四种不同的均衡策略下，先验概率对服务器效用的影响

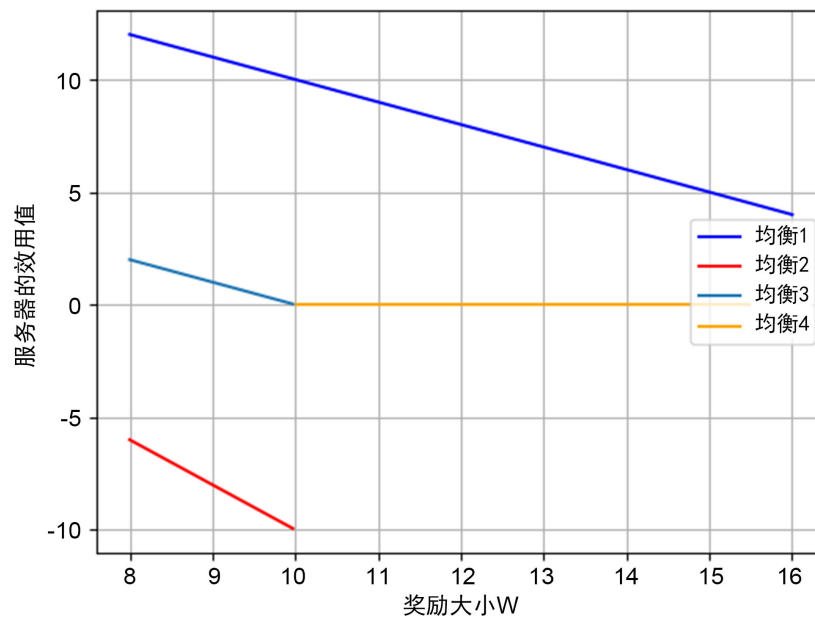
其次，由图 4 观察了在不同均衡条件下，模型价值对服务器效用的影响。图 4(a)中，模型价值的增加对服务器效用是非递减的，并且在均衡 1 策略下，服务器的效用最高，并且模型价值每增加 1 单位，服务器效用的增值最大；图 4(b)中，服务器付出的报酬越多，对服务器效用的影响非递减的，在均衡 2 的策略下，服务器的效用最低，且报酬每增加 1 单位，服务器效用的减值最大，由此可见，对于服务器来说，均衡 1 的策略为最优委托策略，而均衡 2 策略下，服务器的收益最低。

最后，由图 5 观察了在不同均衡条件下，用户所获得的报酬对不同类型用户效用的影响。图 5 中，在断点之前，对于两种类型的用户来说，均衡 2 策略下效用最高，即两种用户都选择{低质量训练}，服务器选择{委托}的策略，而对于服务器来说，均衡 2 的策略为低收益策略，服务器不会选择该均衡策略，故本文所用的博弈方法防止了用户都选择{低质量训练}，服务器选择{委托}的情况发生。



(a) The impact of model value V on server utility

(a) 模型价值 V 对服务器效用的影响

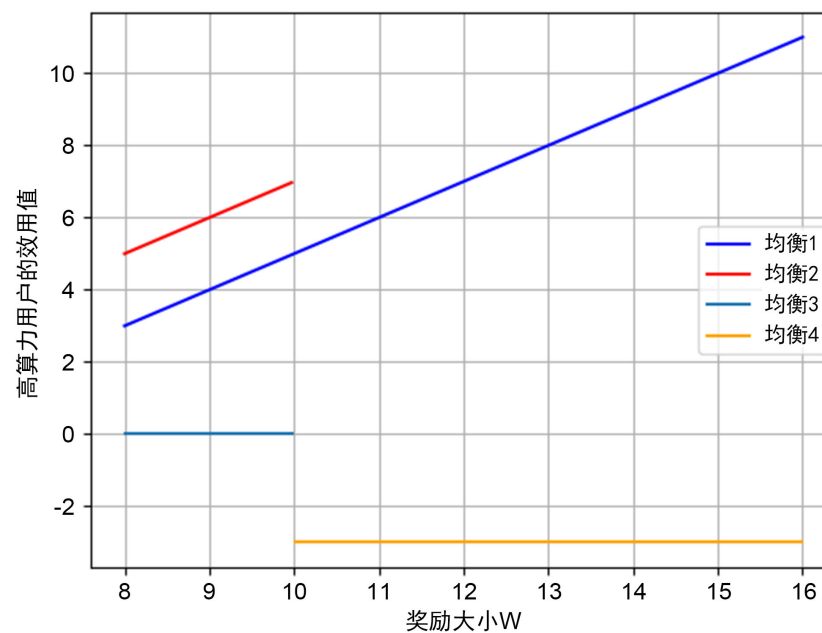


(b) The impact of reward W on server utility
(b) 报酬 W 对服务器效用的影响

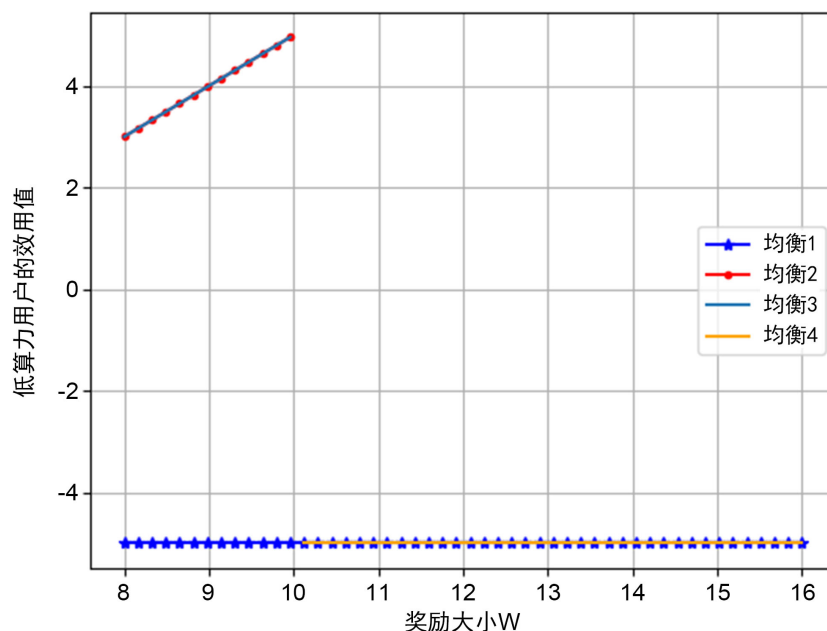
Figure 4. The influence of model value and reward on server utility under different equilibrium strategies

图 4. 不同的均衡策略下模型价值与报酬对服务器效用的影响

由此可知，在所构建的联邦学习博弈模型下，服务器的最优委托策略是委托那些诚实训练的用户，这迫使高计算能力和低计算能力的用户都尽其努力诚实训练。当用户错误行动，选择不诚实训练时，则服务器一定不委托该用户训练任务。



(a) The impact of reward W on the utility of high-power users
(a) 报酬 W 对高算力用户效用的影响

(b) The impact of reward W on the utility of low-power users(b) 报酬 W 对低算力用户效用的影响**Figure 5.** The influence of reward on the utility of two types of users under different equilibrium strategies**图 5.** 不同的均衡策略下报酬对两种类型用户效用的影响

7. 总结

本文研究以信号博弈模型为理论框架，对涉及信息不对称特征的联邦学习场景中各参与者的相互作用进行深入建模。在所构建的模型中，将用户分为高计算能力和低计算能力两类，并通过求解博弈模型的序贯均衡和实验仿真，推导服务器在面对不同类型用户之间的最优委托策略。然而，在研究过程中，我们注意到对模型价值和报酬公平分配的评估方面存在局限。因此，未来的研究将在此基础上展开，重点关注网络贸易交易场景下联邦学习的效用设计，以更全面地考察模型的实际应用和可能的改进方向。

参考文献

- [1] Mammen, P.M. (2021) Federated Learning: Opportunities and Challenges. 1-5. <https://doi.org/10.48550/arXiv.2101.05428>
- [2] Heiss, J., Grünewald, E., Tai, S., et al. (2022) Advancing Blockchain-Based Federated Learning through Verifiable Off-chain Computations. 2022 *IEEE International Conference on Blockchain*, Espoo, 22-25 August 2022, 194-201. <https://doi.org/10.1109/Blockchain55522.2022.00034>
- [3] Data.Konečný, J., McMahan, H.B., Ramage, D., et al. (2016) Federated Optimization: Distributed Machine Learning for On-Device Intelligence. 1-38. <https://doi.org/10.48550/arXiv.1610.02527>
- [4] Banabilah, S., Aloqaily, M., Alsayed, E., et al. (2022) Federated Learning Review: Fundamentals, Enabling Technologies, and Future Applications. *Information Processing & Management*, **59**, Article ID: 103061. <https://doi.org/10.1016/j.ipm.2022.103061>
- [5] Zhang, K., Song, X., Zhang, C., et al. (2022) Challenges and Future Directions of Secure Federated Learning: A Survey. *Frontiers of Computer Science*, **16**, 1-8. <https://doi.org/10.1007/s11704-021-0598-z>
- [6] Liu, J., Huang, J., Zhou, Y., et al. (2022) From Distributed Machine Learning to Federated Learning: A Survey. *Knowledge and Information Systems*, **64**, 885-917. <https://doi.org/10.1007/s10115-022-01664-x>
- [7] Nguyen, D.C., Pham, Q.V., Pathirana, P.N., et al. (2022) Federated Learning for Smart Healthcare: A Survey. *ACM Computing Surveys*, **55**, 1-37. <https://doi.org/10.1145/3501296>
- [8] Tang, M. and Wong, V.W.S. (2021) An Incentive Mechanism for Cross-Silo Federated Learning: A Public Goods

-
- Perspective. *IEEE INFOCOM 2021-IEEE Conference on Computer Communications*, Vancouver, BC, 10-13 May 2021, 1-10. <https://doi.org/10.1109/INFOCOM42981.2021.9488705>
- [9] Pei, J., Zhong, K., Jan, M.A., *et al.* (2022) Personalized Federated Learning Framework for Network Traffic Anomaly Detection. *Computer Networks*, **209**, Article ID: 108906. <https://doi.org/10.1016/j.comnet.2022.108906>
 - [10] Zhang, T., Gao, L., He, C., *et al.* (2022) Federated Learning for the Internet of Things: Applications, Challenges, and Opportunities. *IEEE Internet of Things Magazine*, **5**, 24-29. <https://doi.org/10.1109/IOTM.004.2100182>
 - [11] Wang, Z., Hu, Q., Li, R., *et al.* (2022) Incentive Mechanism Design for Joint Resource Allocation in Blockchain-Based Federated Learning. *IEEE Transactions on Parallel and Distributed Systems*, **34**, 1536-1547.
 - [12] Nishio, T., Shinkuma, R. and Mandayam, N.B. (2020) Estimation of Individual Device Contributions for Incentivizing Federated Learning. 2020 *IEEE Globecom Workshops*, Taipei, 7-11 December 2020, 1-6. <https://doi.org/10.1109/GCWkshps50303.2020.9367484>
 - [13] Li, Q., Wen, Z., Wu, Z., *et al.* (2021) A Survey on Federated Learning Systems: Vision, Hype and Reality for Data Privacy and Protection. *IEEE Transactions on Knowledge and Data Engineering*, **35**, 3347-3366. <https://doi.org/10.1109/TKDE.2021.3124599>
 - [14] Zhang, J., Wu, Y. and Pan, R. (2021) Incentive Mechanism for Horizontal Federated Learning Based on Reputation and Reverse Auction. *Proceedings of the Web Conference*, Ljubljana, April 2021, 947-956. <https://doi.org/10.1145/3442381.3449888>