



(12) 发明专利

(10) 授权公告号 CN 113312169 B

(45) 授权公告日 2023. 12. 19

(21) 申请号 202010123424.2

G06N 20/00 (2019.01)

(22) 申请日 2020.02.27

(56) 对比文件

(65) 同一申请的已公布的文献号

CN 110601814 A, 2019.12.20

申请公布号 CN 113312169 A

CN 110442457 A, 2019.11.12

(43) 申请公布日 2021.08.27

CN 109324902 A, 2019.02.12

(73) 专利权人 香港理工大学深圳研究院

CN 110378488 A, 2019.10.25

地址 518057 广东省深圳市南山区高新园

CN 110263936 A, 2019.09.20

南区粤兴一道18号香港理工大学产学研

CN 110674528 A, 2020.01.10

研大楼205室

US 2019227980 A1, 2019.07.25

审查员 周超群

(72) 发明人 郭嵩 詹玉峰

(74) 专利代理机构 深圳中一专利商标事务所

44237

专利代理师 曹小翠

(51) Int. Cl.

G06F 9/50 (2006.01)

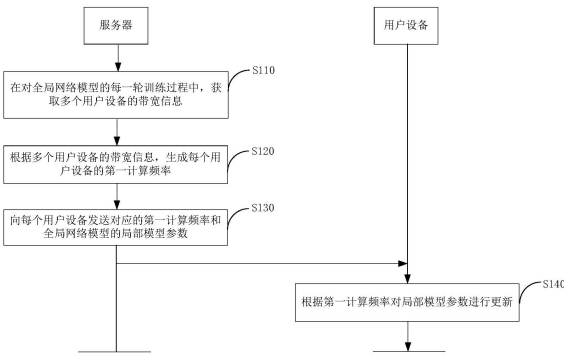
权利要求书2页 说明书10页 附图5页

(54) 发明名称

一种计算资源的分配方法及装置

(57) 摘要

本申请适用于计算机技术领域,提供了一种计算资源的分配方法,包括:在对全局网络模型的每一轮训练过程中,获取多个用户设备的带宽信息,然后根据多个用户设备的带宽信息,生成每个用户设备的第一计算频率,最后向每个用户设备发送对应的第一计算频率和全局网络模型的局部模型参数,以使用户设备根据第一计算频率对局部模型参数进行更新。通过根据多个用户设备的带宽信息生成每个用户设备的第一计算频率,进而每个用户设备可以按照对应的第一计算频率对局部模型参数进行更新,可以使得原先训练任务完成过快的用户设备降低计算速度,从而解决现有联邦学习过程中计算资源浪费的问题。



1. 一种计算资源的分配方法,应用于服务器,其特征在于,所述方法包括:
在对全局网络模型的每一轮训练过程中,获取多个用户设备的带宽信息;
根据所述多个用户设备的带宽信息,生成每个所述用户设备的第一计算频率;
向每个所述用户设备发送对应的第一计算频率和所述全局网络模型的局部模型参数,
以使所述用户设备根据所述第一计算频率对所述局部模型参数进行更新。

2. 根据权利要求1所述的计算资源的分配方法,其特征在于,所述根据所述多个用户设备的带宽信息,生成每个所述用户设备的第一计算频率,包括:

将所述多个用户设备的带宽信息输入到已训练的强化模型中处理,输出每个所述用户设备的第一计算频率,其中,所述强化模型基于训练样本集合中的多个带宽样本训练得到,所述带宽样本包括所述多个用户设备的历史带宽信息。

3. 根据权利要求2所述的计算资源的分配方法,其特征在于,所述将所述多个用户设备的带宽信息输入到已训练的强化模型中处理之前,所述方法还包括:

将所述训练样本集合中的所述带宽样本输入到初始强化模型中处理,输出每个所述用户设备的第二计算频率;

利用每个所述用户设备的第二计算频率对所述全局网络模型进行训练,得到本轮训练的经验数据,所述经验数据包括完成所述本轮训练的总时间,每个所述用户设备在本轮训练中的功耗、对应的带宽样本,以及每个所述用户设备下一轮训练中对应的带宽样本;

当累计获得N轮训练的经验数据时,对所述N轮训练的经验数据进行处理,得到损失值,其中,N为大于1的整数;

当所述损失值不满足预设条件时,调整所述初始强化模型的模型参数,并返回执行将所述训练样本集合中的所述带宽样本输入初始强化模型中处理,输出每个所述用户设备的第二计算频率的步骤;

当所述损失值满足所述预设条件时,停止训练所述初始强化模型,并将训练后的所述初始强化模型作为所述强化模型。

4. 根据权利要求3所述的计算资源的分配方法,其特征在于,所述训练样本集中还包括每个所述用户设备的最高计算频率,所述将所述训练样本集合中的所述带宽样本输入到初始强化模型中处理,输出每个所述用户设备的第二计算频率,包括:

将所述训练样本集合中的所述带宽样本输入到所述初始强化模型中处理,得到每个所述用户设备的第三计算频率;

若所述第三计算频率大于或等于对应的所述最高计算频率,则将所述最高计算频率作为所述第二计算频率输出;

若所述第三计算频率小于对应的所述最高计算频率,则将所述第三计算频率作为所述第二计算频率输出。

5. 根据权利要求3所述的计算资源的分配方法,其特征在于,所述利用每个所述用户设备的第二计算频率对全局网络模型进行训练,得到本轮训练的经验数据,包括:

向每个所述用户设备发送对应的第二计算频率和所述全局网络模型的局部模型参数,以使所述用户设备根据所述第二计算频率对所述局部模型参数进行更新;

接收每个所述用户设备发送的更新后的局部模型参数和所述功耗信息,所述功耗信息为所述用户设备根据对应的第二计算频率对对应的局部模型参数进行更新所产生的功耗;

将每个所述用户设备的更新后的局部模型参数进行组合,得到所述全局网络模型更新后的全局模型参数;

统计完成所述本轮训练的总时间。

6.一种计算资源的分配方法,应用于用户设备,其特征在于,所述方法包括:

接收服务器发送的第一计算频率和全局网络模型的局部模型参数,所述第一计算频率为所述服务器根据多个所述用户设备的带宽信息生成的;

根据所述第一计算频率对所述局部模型参数进行更新。

7.根据权利要求6所述的计算资源的分配方法,其特征在于,所述方法还包括:

接收所述服务器发送的第二计算频率和所述全局网络模型的局部模型参数,所述第二计算频率为所述服务器将训练样本集中的所述带宽样本输入到初始强化模型中处理得到的;

根据所述第二计算频率对所述局部模型参数进行更新;

统计功耗信息,并向所述服务器发送更新后的局部模型参数和所述功耗信息,所述功耗信息为所述用户设备根据对应的第二计算频率对对应的局部模型参数进行更新所产生的功耗。

8.一种计算资源的分配装置,其特征在于,所述装置包括:

获取模块,用于在对全局网络模型的每一轮训练过程中,获取多个用户设备的带宽信息;

生成模块,用于根据所述多个用户设备的带宽信息,生成每个所述用户设备的第一计算频率;

发送模块,用于向每个所述用户设备发送对应的第一计算频率和所述全局网络模型的局部模型参数,以使所述用户设备根据所述第一计算频率对所述局部模型参数进行更新。

9.一种计算资源的分配装置,其特征在于,所述装置包括:

接收模块,用于接收服务器发送的第一计算频率和全局网络模型的局部模型参数,所述第一计算频率为所述服务器根据多个所述用户设备的带宽信息生成的;

更新模块,用于根据所述第一计算频率对所述局部模型参数进行更新。

10.一种计算机可读存储介质,所述计算机可读存储介质存储有计算机程序,其特征在于,所述计算机程序被处理器执行时实现根据权利要求1至7任一项所述的方法。

一种计算资源的分配方法及装置

技术领域

[0001] 本申请属于计算机技术领域,尤其涉及一种计算资源的分配方法及装置。

背景技术

[0002] 联邦学习是一种分布式的智能学习框架,即将训练任务分发给各个用户设备,以由各个用户设备利用本地的用户隐私数据进行训练,并将训练结果返回至服务器,以使得服务器可以在不直接获取用户隐私数据的情况下构建学习模型,有效的保护了用户隐私。

[0003] 然而,在训练过程中每个用户设备会随机选择计算频率进行训练,导致有的用户设备先完成训练任务,有的用户设备后完成训练任务。但是服务器要集合本轮所有的训练结果才可以进行下一轮的训练。对于相同的计算量,由于用户设备完成处理的时间越短就会消耗越多的功耗,因此,在训练过程中,训练任务完成快的用户设备会产生不必要的功耗,造成了计算资源浪费的问题。

发明内容

[0004] 有鉴于此,本申请实施例提供了一种计算资源的分配方法及装置,可以解决现有技术中,训练任务完成快的用户设备会产生不必要的功耗,造成了计算资源浪费的问题。

[0005] 第一方面,本申请实施例提供了一种计算资源的分配方法,应用于服务器,方法包括:

[0006] 在对全局网络模型的每一轮训练过程中,获取多个用户设备的带宽信息;

[0007] 根据多个用户设备的带宽信息,生成每个用户设备的第一计算频率;

[0008] 向每个用户设备发送对应的第一计算频率和全局网络模型的局部模型参数,以使用户设备根据第一计算频率对局部模型参数进行更新。

[0009] 可选的,根据多个用户设备的带宽信息,生成每个用户设备的第一计算频率,包括:

[0010] 将多个用户设备的带宽信息输入到已训练的强化模型中处理,输出每个用户设备的第一计算频率,其中,强化模型基于训练样本集中的多个带宽样本训练得到,带宽样本包括多个用户设备的历史带宽信息。

[0011] 可选的,将多个用户设备的带宽信息输入到已训练的强化模型中处理之前,方法还包括:

[0012] 将训练样本集合中的带宽样本输入到初始强化模型中处理,输出每个用户设备的第二计算频率;

[0013] 利用每个用户设备的第二计算频率对全局网络模型进行训练,得到本轮训练的经验数据,经验数据包括完成本轮训练的总时间,每个用户设备在本轮训练中的功耗、对应的带宽样本,以及每个用户设备下一轮训练中对应的带宽样本;

[0014] 当累计获得N轮训练的经验数据时,对N轮训练的经验数据进行处理,得到损失值;

[0015] 当损失值不满足预设条件时,调整初始强化模型的模型参数,并返回执行将训练

样本集合中的所述带宽样本输入初始强化模型中处理,输出每个用户设备的第二计算频率的步骤;

[0016] 当损失值满足预设条件时,停止训练初始强化模型,并将训练后的初始强化模型作为强化模型。

[0017] 可选的,训练样本集中还包括每个用户设备的最高计算频率,将训练样本集合中的所述带宽样本输入到初始强化模型中处理,输出每个用户设备的第二计算频率,包括:

[0018] 将训练样本集合中的所述带宽样本输入到初始强化模型中处理,得到每个用户设备的第三计算频率;

[0019] 若第三计算频率大于或等于对应的最高计算频率,则将最高计算频率作为第二计算频率输出;

[0020] 若第三计算频率小于对应的最高计算频率,则将第三计算频率作为第二计算频率输出。

[0021] 可选的,利用每个用户设备的第二计算频率对全局网络模型进行训练,得到本轮训练的经验数据,包括:

[0022] 向每个用户设备发送对应的第二计算频率和全局网络模型的局部模型参数,以使用户设备根据第二计算频率对局部模型参数进行更新;

[0023] 接收每个用户设备发送的更新后的局部模型参数和功耗信息,功耗信息为用户设备根据对应的第二计算频率对对应的局部模型参数进行更新所产生的功耗;

[0024] 将每个用户设备的更新后的局部模型参数进行组合,得到全局网络模型更新后的全局模型参数;

[0025] 统计完成本轮训练的总时间。

[0026] 第二方面,本申请实施例提供了一种计算资源的分配方法,应用于用户设备,方法包括:

[0027] 接收服务器发送的第一计算频率和全局网络模型的局部模型参数,第一计算频率为服务器根据多个用户设备的带宽信息生成的;

[0028] 根据第一计算频率对局部模型参数进行更新。

[0029] 可选的,方法还包括:

[0030] 接收服务器发送的第二计算频率和全局网络模型的局部模型参数,第二计算频率为服务器将训练样本集合中的所述带宽样本输入到初始强化模型中处理得到的;

[0031] 根据第二计算频率对局部模型参数进行更新;

[0032] 统计功耗信息,并向服务器发送更新后的局部模型参数和功耗信息,功耗信息为用户设备根据对应的第二计算频率对对应的局部模型参数进行更新所产生的功耗。

[0033] 第三方面,本申请实施例提供了一种计算资源的分配装置,装置包括:

[0034] 获取模块,用于在对全局网络模型的每一轮训练过程中,获取多个用户设备的带宽信息;

[0035] 生成模块,用于根据多个用户设备的带宽信息,生成每个用户设备的第一计算频率;

[0036] 发送模块,用于向每个用户设备发送对应的第一计算频率和全局网络模型的局部模型参数,以使用户设备根据第一计算频率对局部模型参数进行更新。

- [0037] 第四方面,本申请实施例提供了一种计算资源的分配装置,装置包括:
- [0038] 接收模块,用于接收服务器发送的第一计算频率和全局网络模型的局部模型参数,第一计算频率为服务器根据多个用户设备的带宽信息生成的;
- [0039] 更新模块,用于根据第一计算频率对局部模型参数进行更新。
- [0040] 第五方面,本申请实施例提供了一种计算机可读存储介质,包括计算机可读存储介质存储有计算机程序,计算机程序被处理器执行时实现上述第一方面或第二方面的任一实施方式的方法。
- [0041] 本申请提供的计算资源的分配方法及装置,通过在对全局网络模型的每一轮训练过程中,获取多个用户设备的带宽信息,然后根据多个用户设备的带宽信息,生成每个用户设备的第一计算频率,最后向每个用户设备发送对应的第一计算频率和全局网络模型的局部模型参数,以使用户设备根据第一计算频率对局部模型参数进行更新。通过根据多个用户设备的带宽信息生成每个用户设备的第一计算频率,进而每个用户设备可以按照对应的第一计算频率对局部模型参数进行更新,可以使得原先训练任务完成过快的用户设备降低计算速度,从而解决现有联邦学习过程中计算资源浪费的问题。

附图说明

- [0042] 为了更清楚地说明本申请实施例中的技术方案,下面将对实施例或现有技术描述中所需要使用的附图作简单地介绍,显而易见地,下面描述中的附图仅仅是本申请的一些实施例,对于本领域普通技术人员来讲,在不付出创造性劳动性的前提下,还可以根据这些附图获得其他的附图。
- [0043] 图1是本申请提供的一种联邦学习的示意性流程图;
- [0044] 图2是本申请提供的一种在静态环境下联邦学习的训练时间对比示意图;
- [0045] 图3是本申请提供的一种在动态环境下联邦学习的训练时间对比示意图;
- [0046] 图4是本申请提供的一种计算资源的分配方法的示意性流程图;
- [0047] 图5是本申请提供的一种强化模型的训练流程图;
- [0048] 图6是本申请提供的一种深度强化学习网络的训练示意图;
- [0049] 图7是本申请提供的一种计算资源的分配装置的结构框图;
- [0050] 图8是本申请提供的另一种计算资源的分配装置的结构框图。

具体实施方式

- [0051] 以下描述中,为了说明而不是为了限定,提出了诸如特定系统结构、技术之类的具体细节,以便透彻理解本申请实施例。然而,本领域的技术人员应当清楚,在没有这些具体细节的其它实施例中也可以实现本申请。在其它情况中,省略对众所周知的系统、装置、电路以及方法的详细说明,以免不必要的细节妨碍本申请的描述。
- [0052] 应当理解,当在本申请说明书和所附权利要求书中使用时,术语“包括”指示所描述特征、整体、步骤、操作、元素和/或组件的存在,但并不排除一个或多个其它特征、整体、步骤、操作、元素、组件和/或其集合的存在或添加。
- [0053] 还应当理解,在本申请说明书和所附权利要求书中使用的术语“和/或”是指相关联列出的项中的一个或多个的任何组合以及所有可能组合,并且包括这些组合。

[0054] 另外,在本申请说明书和所附权利要求书的描述中,术语“第一”、“第二”、“第三”等仅用于区分描述,而不能理解为指示或暗示相对重要性。

[0055] 本申请实施例提供的计算资源的分配方法可以应用于手机、平板电脑、可穿戴设备、车载设备、服务器、笔记本电脑、超级移动个人计算机(ultra-mobile personal computer,UMPC)、上网本、个人数字助理(personal digital assistant,PDA)等计算设备上。本申请实施例对计算设备的具体类型不作任何限制。

[0056] 联邦学习是一种分布式的智能学习框架,能有效帮助企业在满足用户隐私保护、数据安全和政府法规的要求下,进行数据使用和机器学习建模。图1示出了本申请提供的一种联邦学习的示意性流程图。在联邦学习的场景中,通常服务器需要将待构建的全局网络模型转化为多个局部模型参数,然后将多个局部模型参数通过不同的网络通道分发到对应的用户设备中。其中,用户设备可以是移动终端,或计算机等电子设备,服务器可以通过无线网络通道向移动设备发送局部模型参数,也可以通过有线网络通道向计算机发送局部模型参数。用户设备在接收到局部模型参数后,可以在本地根据本地的用户数据进行更新,并将更新后的局部模型参数返回至服务器。服务器可以根据多个更新后的局部模型参数进行训练,并在经过多轮的训练后判断全局网络模型是否满足训练结束条件。

[0057] 在联邦学习的过程中,每个用户设备都可能是不一样的,且每个用户设备都可以根据本地的预设规则或随机选择计算频率来更新局部模型参数。这就导致有的用户设备先完成训练任务,有的用户设备后完成训练任务。而服务器必须等接收到本轮所有的更新后的局部模型参数才可以进行下一轮的训练。对于相同的计算量,由于用户设备完成处理的时间越短就会消耗越多的功耗,因此,训练任务完成快的用户设备并不能提高整体训练速度,反而会产生不必要的功耗,浪费了计算资源。

[0058] 例如,图2示出了本申请提供的一种在静态环境下联邦学习的训练时间对比示意图。本轮有3个用户设备参与训练,每个用户设备的带宽都是一样的,因此每个用户设备传输数据的时间也是一样的。在本地计算数据方面,用户设备1的计算频率最低,因此训练任务完成的最慢,花费的时间最多;用户设备2的计算频率最高,因此训练任务完成的最快,花费的时间最少;用户设备3的计算频率适中,因此训练任务花费的时间适中。可以看出本轮训练完成时间是以用户设备1花费的时间确定的,所以训练任务完成快的用户设备2和用户设备3并不能提高整体训练速度,反而浪费了不必要的功耗,造成了计算资源浪费的问题。

[0059] 进一步的,在实际应用中,每个用户设备的带宽都可能是不一样的,且每轮训练时用户设备选择的计算频率也可能是不一样的。图3示出了本申请提供的一种在动态环境下联邦学习的训练时间对比示意图。可以看出每轮训练时,每个用户设备的带宽和计算频率都不一样的,因此,在动态环境中上述问题变得更加复杂。

[0060] 针对上述问题,本申请提供了一种计算资源的分配方法,图4示出了本申请提供的一种计算资源的分配方法的示意性流程图。

[0061] S110、在对全局网络模型的每一轮训练过程中,获取多个用户设备的带宽信息。

[0062] 在每轮训练中,总时间由用户设备的计算时间和数据传输的传输时间组成。因此,我们可以通过控制用户设备的计算频率来控制计算时间,并在用户设备一完成本地训练时就上传数据,进而控制用户设备的上传时刻。由于上传时刻是可控的,所以我们只要判断出数据传输阶段的带宽信息就可以得到传输时间,最终得到每个用户设备的总时间。在可以

获取到每个用户设备的总时间的情况下,通过调整每个用户设备的计算频率就可以实现服务器可以在同一时刻接收到每个用户设备发送的数据,同时每个用户设备也没有产生多余的功耗。

[0063] 需要说明的是,本申请实施例中的计算频率在实际应用中,可以使用中央处理器(Central Processing Unit,CPU)的计算频率作为实际值。

[0064] 在一个实施例中,服务器可以在对全局网络模型的每一轮训练过程中,在下发局部模型参数前获取全部参与训练的用户设备的带宽信息。在实际应用中,网络链路的带宽可以看作是无规律变化的,但根据现有的研究表明,在短时间内动态变化的带宽可以看作是静态不变的,所以服务器可以选取一段很短的时间内的带宽。例如,服务器可以选取当前时刻的前10秒内的带宽信息,作为每个用户设备的带宽信息。

[0065] S120、根据多个用户设备的带宽信息,生成每个用户设备的第一计算频率。

[0066] 服务器可以根据每个用户设备的带宽信息,预估每个用户设备在传输数据时的传输时间,再获取每个用户设备的计算频率,计算每个用户设备在本地训练需要花费的计算时间。服务器可以通过改变每个用户设备的计算频率,调整每个用户设备完成本轮训练的总时间,直到每个用户设备完成本轮训练的总时间是一样的。服务器就可以确定此时每个用户设备对应的计算频率。

[0067] 例如,在每轮训练中包括,用户设备1、用户设备2和用户设备3。用户设备1的带宽中等大小,用户设备2的带宽最小,用户设备3的带宽最大。首先服务器可以预估出用户设备1的传输时间适中,用户设备2的传输时间最长,用户设备3的传输时间最短。然后,服务器可以获取用户设备2的计算频率,并计算出用户设备2的计算时间,将用户设备2的计算时间与传输时间相加得到总时间。服务器再将用户设备2的总时间分别减去用户设备1的传输时间和用户设备3的传输时间,得到用户设备1的计算时间和用户设备3的计算时间。最后根据用户设备1的计算时间和用户设备3的计算时间计算出用户设备1的第一计算频率和用户设备3的第一计算频率。

[0068] 可选的,服务器中设置有已训练的强化模型。其中,强化模型可以是基于深度强化学习(Deep Reinforcement Learning,DRL)网络训练所得。以使服务器可以将多个用户设备的带宽信息输入到已训练的强化模型中处理,输出每个用户设备的第一计算频率,其中,强化模型基于训练样本集中的多个带宽样本训练得到,带宽样本包括多个用户设备的历史带宽信息。在实际应用中,会有很多用户设备参与到训练的过程中,因此,根据多个用户设备的带宽信息生成每个用户设备的第一计算频率是一个复杂的计算过程。通过DRL网络可以快速、准确的获得计算结果。

[0069] S130、向每个用户设备发送对应的第一计算频率和全局网络模型的局部模型参数,以使用户设备根据第一计算频率对局部模型参数进行更新。

[0070] 服务器在生成每个用户设备对应的第一计算频率的情况下,就可以将每个用户设备对应的第一计算频率和全局网络模型的局部模型参数发送给对应的用户设备,以使用户设备根据第一计算频率对局部模型参数进行更新。

[0071] 进一步的,用户设备可以接收服务器发送的第一计算频率和全局网络模型的局部模型参数,第一计算频率为服务器根据多个用户设备的带宽信息生成的。

[0072] 例如,用户设备1可以接收服务器发送的第一计算频率1和局部模型参数1,用户设

备2可以接收服务器发送的第一计算频率2和局部模型参数2等。

[0073] S140、根据第一计算频率对局部模型参数进行更新。

[0074] 用户设备可以将本地的计算频率调整至第一计算频率,并根据第一计算频率对局部模型参数进行更新。其中,更新是利用本地数据对局部模型参数进行训练,得到更新后的局部模型参数的过程。在生成更新后的局部模型参数的情况下,用户设备可以将更新后的局部模型参数发送至服务器,以供服务器构建满足条件的全局网络模型。

[0075] 在本发明实施例中,服务器可以根据多个用户设备的带宽信息生成每个用户设备的第一计算频率,并通过控制用户设备的计算频率来控制计算时间,使得用户设备在一完成本地训练时就上传数据,进而控制用户设备的上传时刻。保证服务器可以在同一时刻接收到每个用户设备发送的数据,同时每个用户设备没有产生多余的功耗。

[0076] 可选的,图5为本申请提供的一种强化模型的训练流程图,主要涉及基于训练样本集对强化模型进行训练的过程。参见图5,该训练的方法包括:

[0077] S210、将训练样本集中的带宽样本输入到初始强化模型中处理,输出每个用户设备的第二计算频率。

[0078] 历史带宽信息是用户设备在过去一段时间内的带宽信息。带宽样本是所有用户设备的历史带宽信息的集合。训练样本集是多个带宽样本的集合。服务器可以在离线的环境下根据训练样本集构建一个模拟场景,并在该模拟场景下训练强化模型。

[0079] 在一个实施例中,首先服务器可以初始化强化模型的参数,得到初始强化模型。然后,服务器可以在训练样本集中随机选取一个靠前的时间点作为起始时间,并在此后的训练中根据带宽信息的变化模拟数据传输时的带宽变化。服务器在确定好起始时间后,获取训练样本集中的带宽样本,并将其输入到初始强化模型中处理,输出每个用户设备的第二计算频率。

[0080] 示例性的,在强化模型训练的时候,输出的第二计算频率可能会高于用户设备的最高计算频率,此时,用户设备是无法达到该第二计算频率的。因此,训练样本集中还包括每个用户设备的最高计算频率,步骤S210还可以包括S211-S213:

[0081] S211、将训练样本集中的带宽样本输入到初始强化模型中处理,得到每个用户设备的第三计算频率。

[0082] 服务器在确定好起始时间后,获取训练样本集中的带宽样本,并将其输入到初始强化模型中处理,输出每个用户设备的第三计算频率。

[0083] S212、若第三计算频率大于或等于对应的最高计算频率,则将最高计算频率作为第二计算频率输出。

[0084] 服务器可以将每个用户设备的第三计算频率与对应的最高计算频率对比,若第三计算频率大于或等于对应的最高计算频率,则将最高计算频率作为第二计算频率输出。

[0085] S213、若第三计算频率小于对应的最高计算频率,则将第三计算频率作为第二计算频率输出。

[0086] 服务器可以将每个用户设备的第三计算频率与对应的最高计算频率对比,若第三计算频率小于对应的最高计算频率,则将第三计算频率作为第二计算频率输出。

[0087] 服务器可以根据每个用户设备的最高计算频率,确保每个用户设备都可以按照对应的第二计算频率进行训练。在服务器输出第二计算频率之后,即可执行后面的步骤。

[0088] S220、利用每个用户设备的第二计算频率对全局网络模型进行训练,得到本轮训练的经验数据,经验数据包括完成本轮训练的总时间,每个用户设备在本轮训练中的功耗、对应的带宽样本,以及每个用户设备下一轮训练中对应的带宽样本。

[0089] 服务器可以利用每个用户设备的第二计算频率对全局网络模型进行训练,并在进行训练的同时得到本轮训练的经验数据,其中,经验数据包括完成本轮训练的总时间,每个用户设备的功耗、对应的带宽样本,以及每个用户设备下一轮训练中对应的带宽样本。

[0090] 具体的,步骤S220还可以包括S221-S225:

[0091] S221、向每个用户设备发送对应的第二计算频率和全局网络模型的局部模型参数,以使用户设备根据第二计算频率对局部模型参数进行更新。

[0092] 服务器可以向每个用户设备发送对应的第二计算频率和全局网络模型的局部模型参数。

[0093] 进一步的,用户设备可以接收服务器发送的第二计算频率和全局网络模型的局部模型参数,第二计算频率为服务器将训练样本集合中的带宽样本输入到初始强化模型中处理得到的。

[0094] S222、根据第二计算频率对局部模型参数进行更新;

[0095] 用户设备可以将本地的计算频率调整至第二计算频率,并根据第二计算频率对局部模型参数进行更新。

[0096] S223、统计功耗信息,并向服务器发送更新后的局部模型参数和功耗信息,功耗信息为用户设备根据对应的第二计算频率对对应的局部模型参数进行更新所产生的功耗。

[0097] 用户设备在更新局部模型参数的同时,可以统计本次更新局部模型参数时的功耗信息,其中,功耗信息为用户设备根据对应的第二计算频率对对应的局部模型参数进行更新所产生的功耗。然后,用户设备可以将更新后的局部模型参数和功耗信息发送至服务器。因为训练是在离线状态下进行的,此时数据传输链路可以采用训练样本集合中的历史带宽信息作为数据传输时的带宽状态。

[0098] 进一步的,服务器可以接收每个用户设备发送的更新后的局部模型参数和功耗信息。

[0099] S224、将每个用户设备的更新后的局部模型参数进行组合,得到全局网络模型更新后的全局模型参数。

[0100] 服务器可以将每个用户设备的更新后的局部模型参数进行组合,得到全局网络模型更新后的全局模型参数。

[0101] S225、统计完成本轮训练的总时间。

[0102] 服务器可以统计完成本轮训练的总时间,总时间主要包括两个部分,一个是用户设备在本地对局部模型参数进行更新时所用的计算时间,一个是更新后的局部模型参数在传输过程中所用的传输时间。至此,服务器就得到了本轮训练的经验数据,其中,经验数据包括完成本轮训练的总时间,每个用户设备的功耗、对应的带宽样本,以及每个用户设备下一轮训练中对应的带宽样本。

[0103] 在服务器得到本轮训练的总时间,每个用户设备在本轮训练中的功耗、对应的带宽样本,以及每个用户设备下一轮训练中对应的带宽样本之后,即可执行后面的步骤。

[0104] S230、当累计获得N轮训练的经验数据时,对N轮训练的经验数据进行处理,得到损

失值。

[0105] 服务器可以将得到经验数据保存在数据库中,该数据库中的经验数据集合可以称为经验池。在经过N轮训练后,经验池中累计N轮经验数据,此时服务器可以使用损失函数对经验池中的经验数据进行处理,得到一个损失值。其中,N可以根据实际情况进行选择,例如,N为20,则服务器在经过20轮训练后,将这20轮的经验数据输入到预设的损失函数中计算,生成损失值。

[0106] 其中,损失函数可以是L1范数损失函数、均方误差损失函数、交叉熵损失函数等,对此本申请不做限制。

[0107] S240、当损失值不满足预设条件时,调整初始强化模型的模型参数,并返回执行将训练样本集合中的带宽样本输入初始强化模型中处理,输出每个用户设备的第二计算频率的步骤。

[0108] 服务器可以根据预设条件判断损失值。例如,服务器可以将预设条件设定为损失值是否趋近于0。当损失值为5时,服务器可以判断损失值不满足预设条件。当判断损失值不满足预设条件时,服务器可以调整初始强化模型的模型参数,也即是对初始强化模型进行一次更新。在进行一次更新后,服务器可以将经验池清空,并返回步骤S210。服务器可以根据更新后的初始强化模型生成新的第二计算频率,重新累计经验池,直到再次完成N轮训练时,判断损失值是否满足预设条件。

[0109] 在一个实施例,我们可以采用一种近端策略优化(Proximal Policy Optimization, PPO)的策略梯度方法对初始强化模型进行更新。

[0110] S240、当损失值满足预设条件时,停止训练初始强化模型,并将训练后的初始强化模型作为强化模型。

[0111] 初始强化模型的模型参数在经过多次调整后,会使得损失值满足预设条件,此时服务器就可以停止训练初始强化模型,并将训练后的初始强化模型作为强化模型用于正式的联邦学习中。例如,当损失值为0.5时,服务器可以判断损失值满足预设条件。

[0112] 例如,参照图6,图6为本申请提供的一种深度强化学习网络的训练示意图。在服务器中,包括一个DRL网络,该DRL网络包括动作(actor)网络和评价(critic)网络。动作网络用于根据状态(state)信息也即是带宽信息,输出动作(action)信息也即是第二计算频率。评价网络用于在更新时根据经验(reward)信息也即是经验数据更新动作网络。经验池用于保存每轮训练中的经验数据。其中,动作网络是强化模型,评价网络中包括预设的损失(loss)函数,用于动作网络得到的经验数据进行处理得到损失值。在开始训练DRL网络前,服务器可以先将DRL网络的参数初始化。然后将带宽信息输入到动作网络中,并输出第二计算频率。服务器将每个用户设备对应的第二计算频率和局部模型参数发送到对应的用户设备。用户设备再将更新后的局部模型参数和功耗信息返回至服务器。服务器可以根据局部模型参数和功耗信息,生成经验数据和新的局部模型参数。服务器在经过N轮训练后,通过损失函数生成损失值,通过损失值判断动作网络是否收敛,若没有收敛则根据经验池对动作网络进行一次更新。再次重复训练步骤并累计经验数据,直到动作网络收敛,停止对动作网络的训练,此时的动作网络就是强化模型。

[0113] 本申请提供的计算资源的分配方法,通过在对全局网络模型的每一轮训练过程中,获取多个用户设备的带宽信息,然后根据多个用户设备的带宽信息,生成每个用户设备

的第一计算频率,最后向每个用户设备发送对应的第一计算频率和全局网络模型的局部模型参数,以使用户设备根据第一计算频率对局部模型参数进行更新。通过根据多个用户设备的带宽信息生成每个用户设备的第一计算频率,进而每个用户设备可以按照对应的第一计算频率对局部模型参数进行更新,可以使得原先训练任务完成过快的用户设备降低计算速度,从而解决现有联邦学习过程中计算资源浪费的问题。

[0114] 应理解,上述实施例中各步骤的序号的大小并不意味着执行顺序的先后,各过程的执行顺序应以其功能和内在逻辑确定,而不对本申请实施例的实施过程构成任何限定。

[0115] 对应于上文实施例所述的计算资源的分配方法,图7示出了本申请提供的一种计算资源的分配装置的结构框图,为了便于说明,仅示出了与本申请实施例相关的部分。

[0116] 参照图7,该装置包括:

[0117] 获取模块710,用于在对全局网络模型的每一轮训练过程中,获取多个用户设备的带宽信息;

[0118] 生成模块720,用于根据多个用户设备的带宽信息,生成每个用户设备的第一计算频率;

[0119] 发送模块730,用于向每个用户设备发送对应的第一计算频率和全局网络模型的局部模型参数,以使用户设备根据第一计算频率对局部模型参数进行更新。

[0120] 图8示出了本申请提供的另一种计算资源的分配装置的结构框图,为了便于说明,仅示出了与本申请实施例相关的部分。

[0121] 参照图8,该装置包括:

[0122] 接收模块810,用于接收服务器发送的第一计算频率和全局网络模型的局部模型参数,第一计算频率为服务器根据多个用户设备的带宽信息生成的;

[0123] 更新模块820,用于根据第一计算频率对局部模型参数进行更新。

[0124] 本申请实施例还提供了一种计算机可读存储介质,所述计算机可读存储介质存储有计算机程序,所述计算机程序被处理器执行时实现可实现上述各个方法实施例中的步骤。

[0125] 需要说明的是,上述装置/单元之间的信息交互、执行过程等内容,由于与本申请方法实施例基于同一构思,其具体功能及带来的技术效果,具体可参见方法实施例部分,此处不再赘述。

[0126] 所属领域的技术人员可以清楚地了解到,为了描述的方便和简洁,仅以上述各功能单元、模块的划分进行举例说明,实际应用中,可以根据需要而将上述功能分配由不同的功能单元、模块完成,即将所述装置的内部结构划分成不同的功能单元或模块,以完成以上描述的全部或者部分功能。实施例中的各功能单元、模块可以集成在一个处理单元中,也可以是各个单元单独物理存在,也可以两个或两个以上单元集成在一个单元中,上述集成的单元既可以采用硬件的形式实现,也可以采用软件功能单元的形式实现。另外,各功能单元、模块的具体名称也只是为了便于相互区分,并不用于限制本申请的保护范围。上述系统中单元、模块的具体工作过程,可以参考前述方法实施例中的对应过程,在此不再赘述。

[0127] 在上述实施例中,对各个实施例的描述都各有侧重,某个实施例中未详述或记载的部分,可以参见其它实施例的相关描述。

[0128] 本领域普通技术人员可以意识到,结合本文中所公开的实施例描述的各示例的单元及算法步骤,能够以电子硬件、或者计算机软件和电子硬件的结合来实现。这些功能究竟以硬件还是软件方式来执行,取决于技术方案的特定应用和设计约束条件。专业技术人员可以对每个特定的应用来使用不同方法来实现所描述的功能,但是这种实现不应认为超出本申请的范围。

[0129] 以上所述实施例仅用以说明本申请的技术方案,而非对其限制;尽管参照前述实施例对本申请进行了详细的说明,本领域的普通技术人员应当理解:其依然可以对前述各实施例所记载的技术方案进行修改,或者对其中部分技术特征进行等同替换;而这些修改或者替换,并不使相应技术方案的本质脱离本申请各实施例技术方案的精神和范围,均应包含在本申请的保护范围之内。

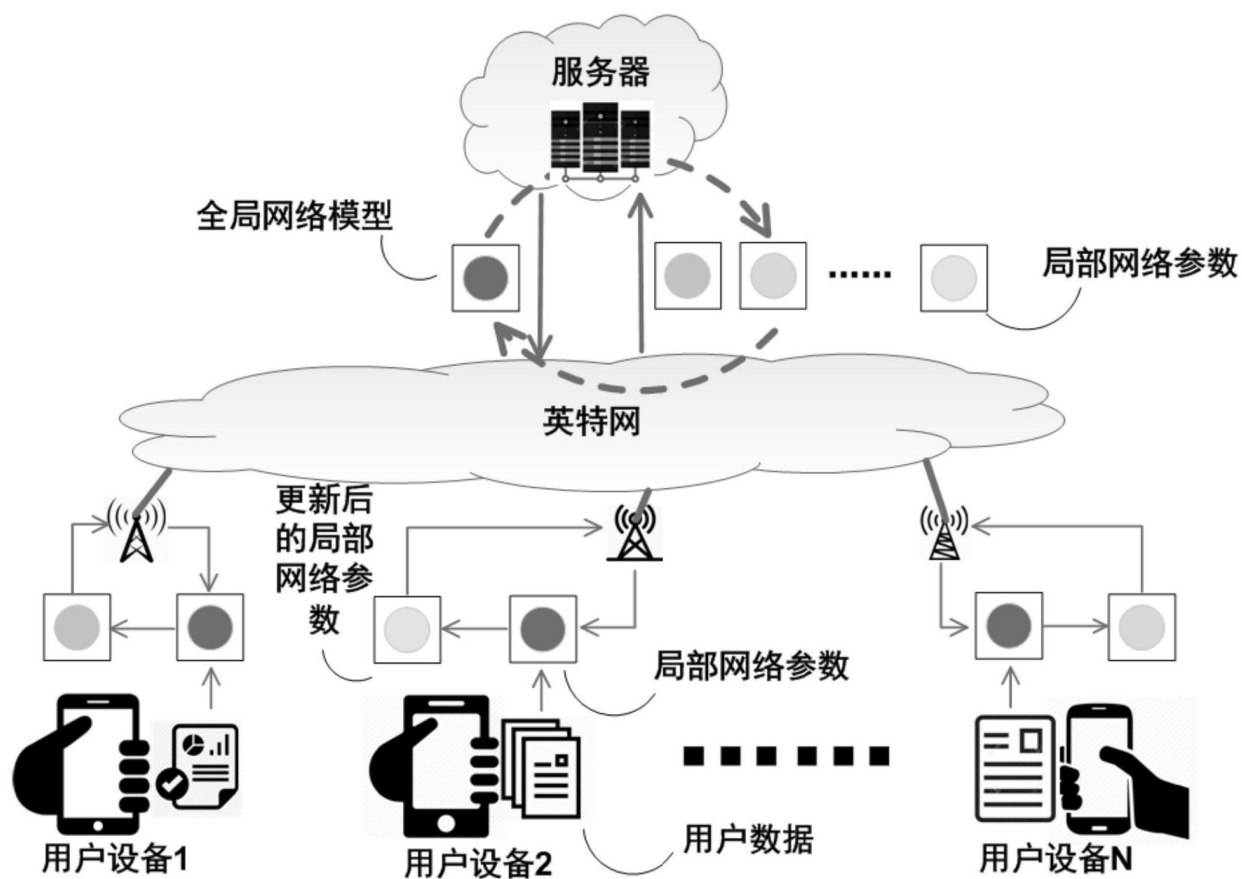


图1

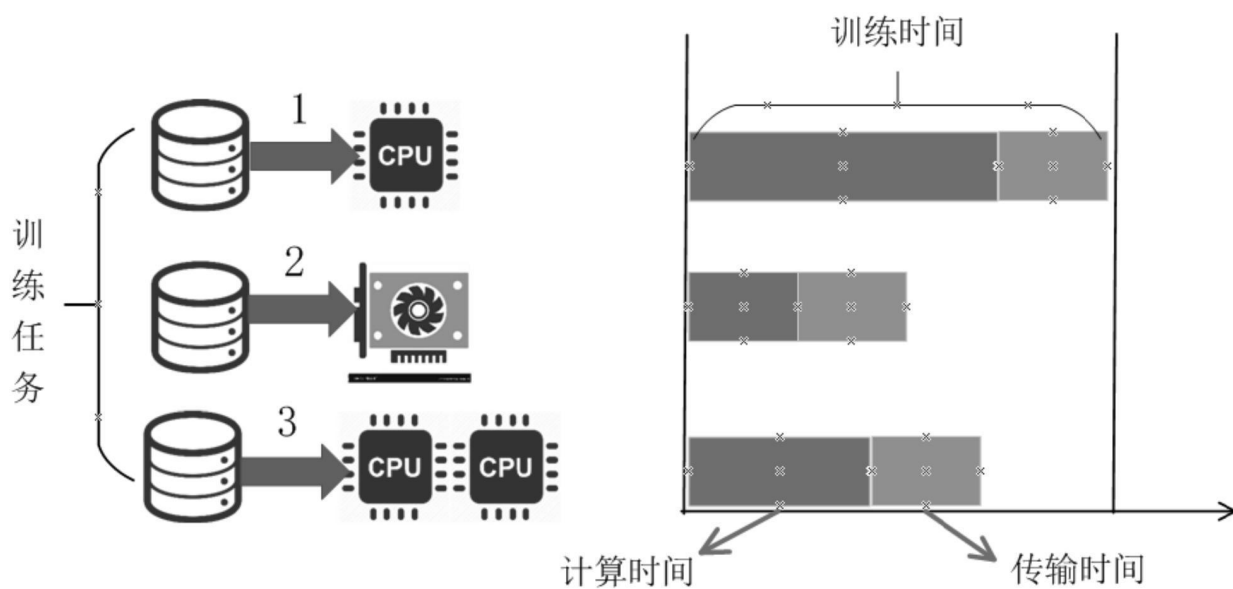


图2

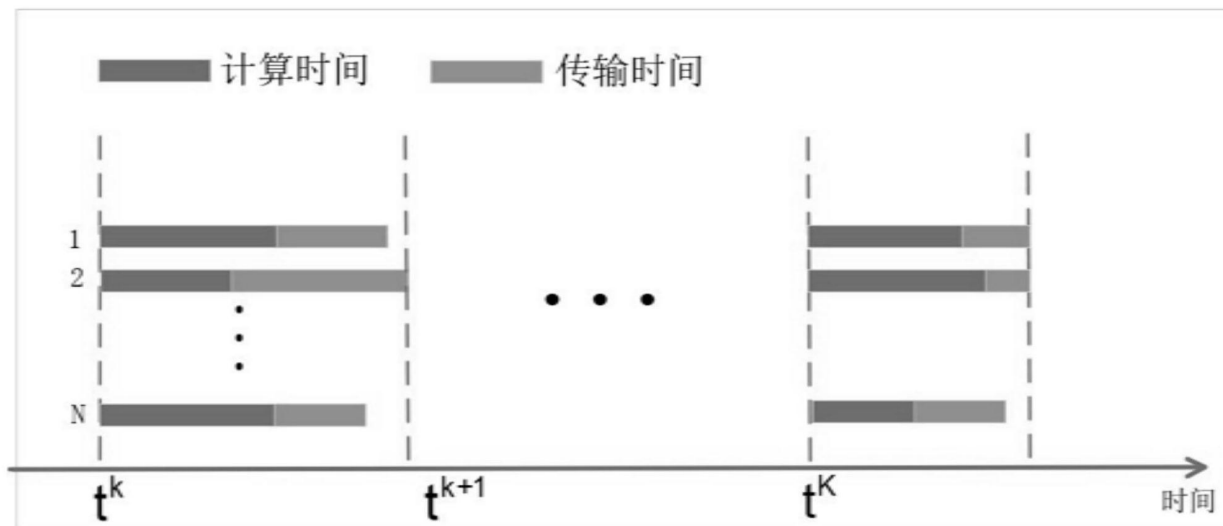


图3

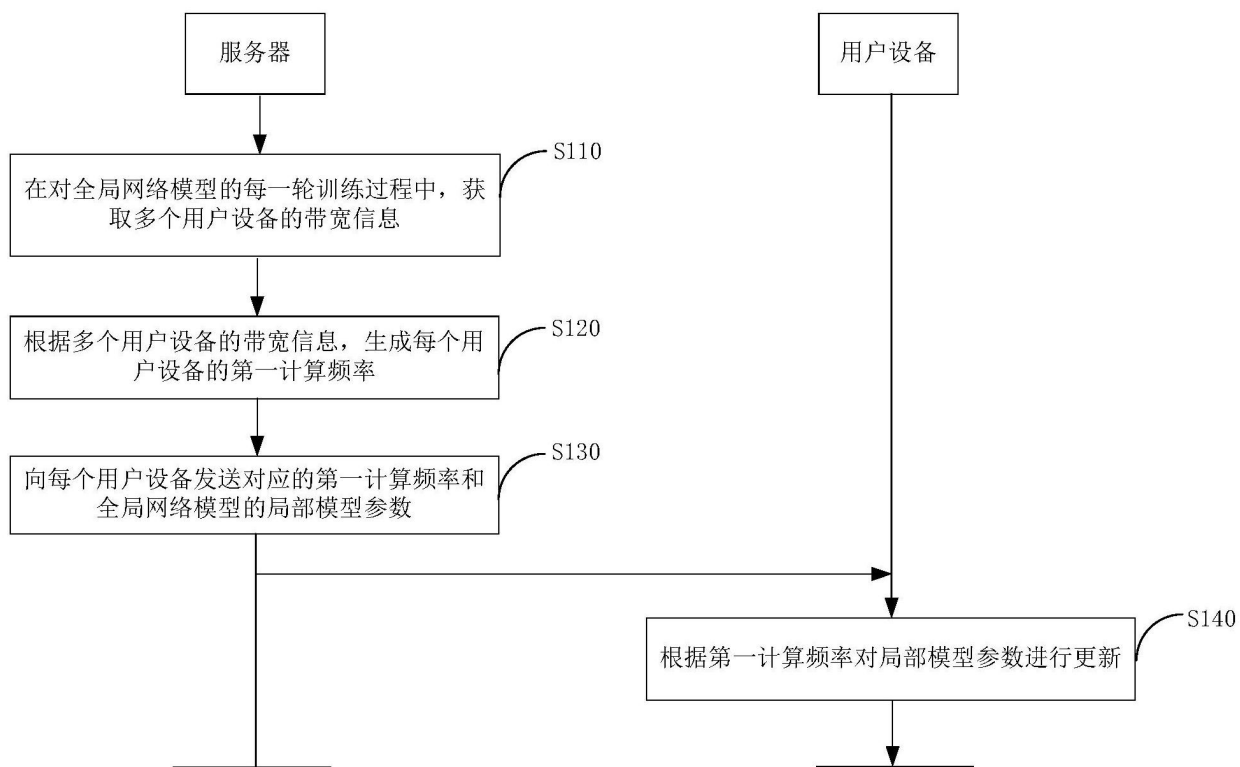


图4

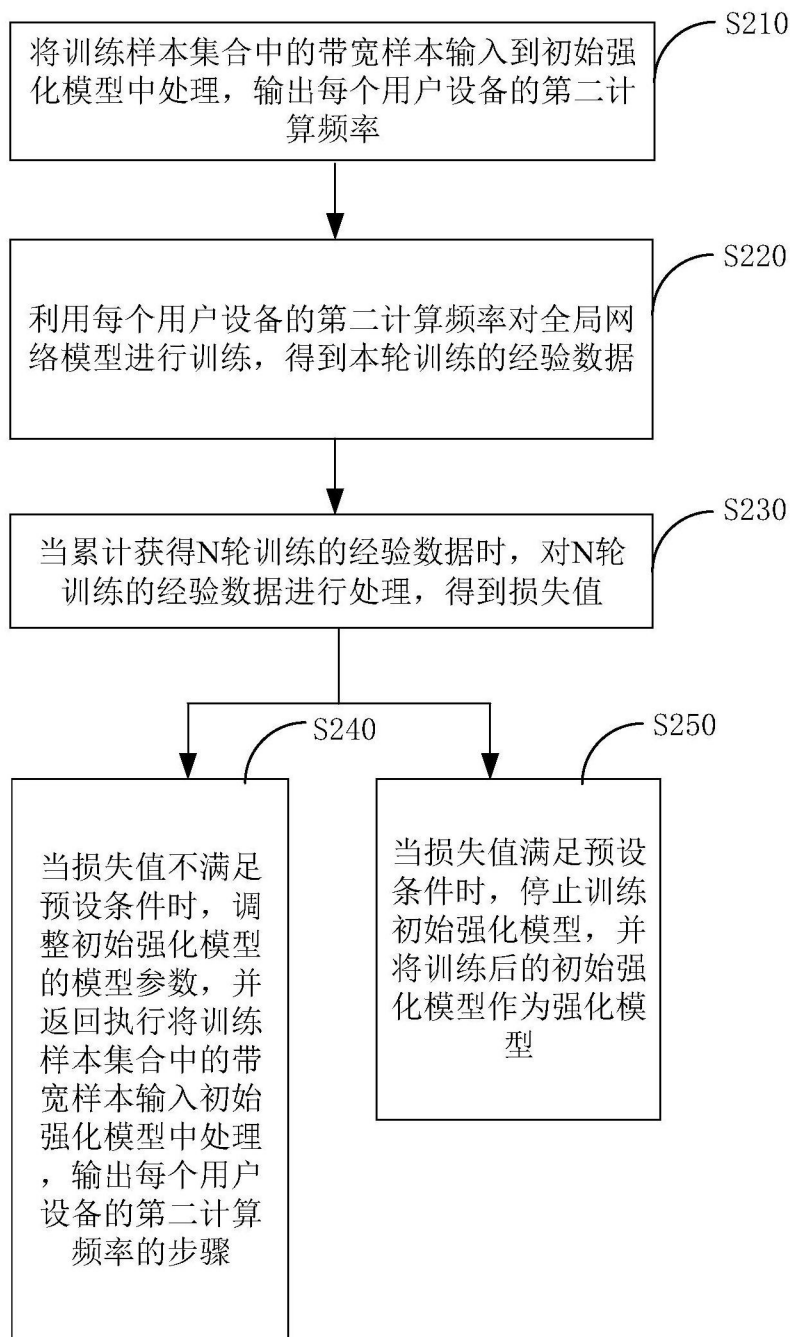


图5

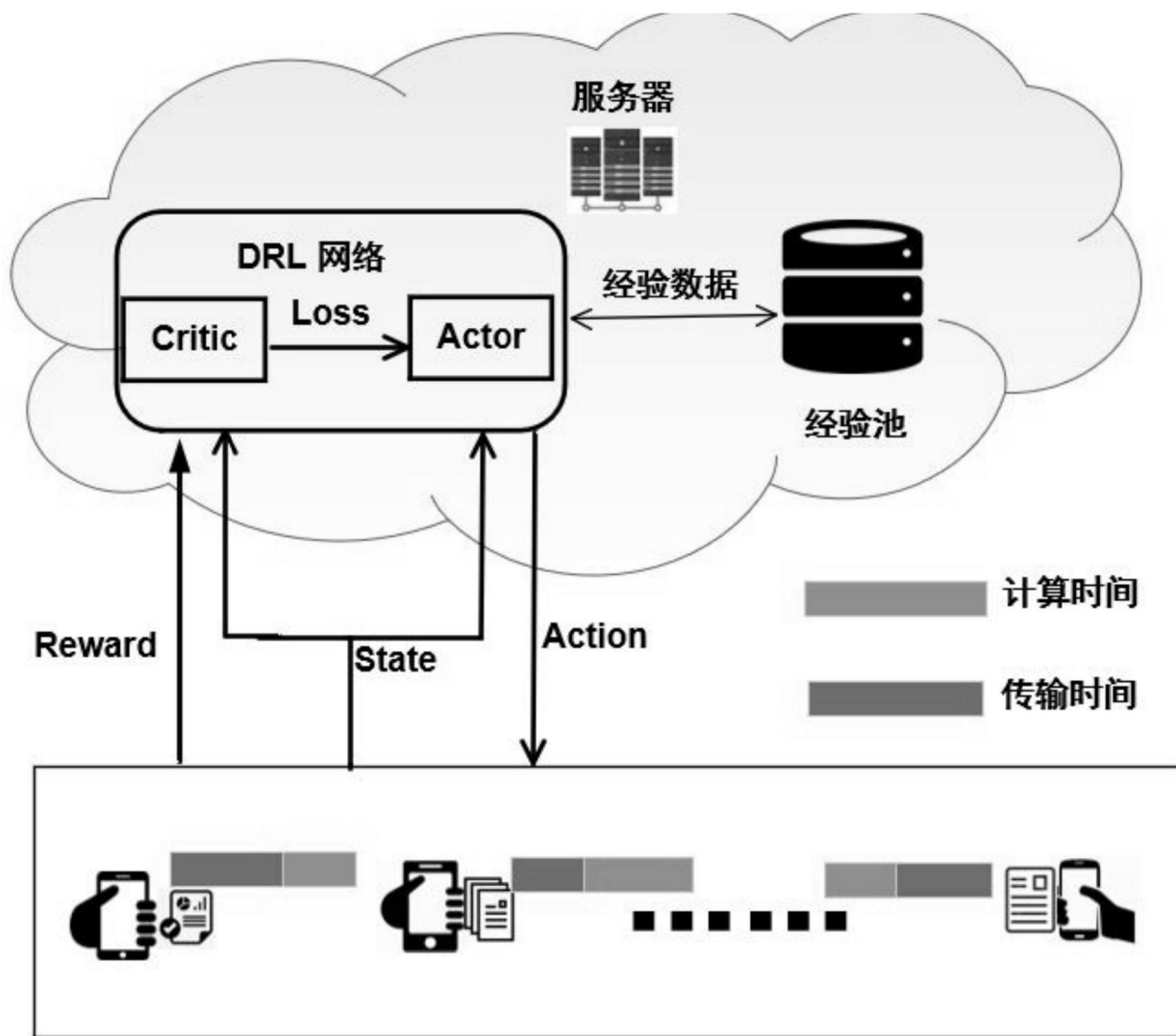


图6

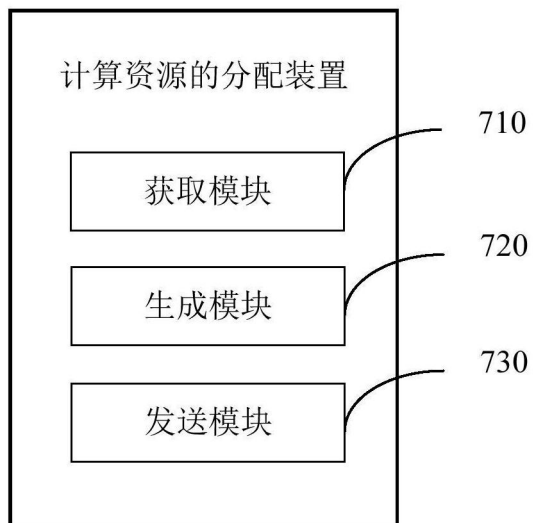


图7

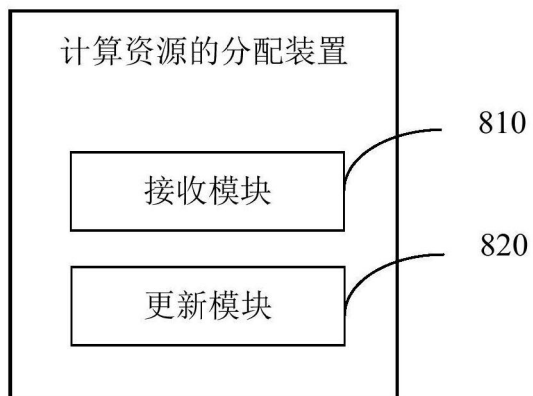


图8