

Computation Energy Efficiency Maximization for Intelligent Reflective Surface-Aided Wireless Powered Mobile Edge Computing

Junhui Du, Minxian Xu, *Member, IEEE*, Sukhpal Singh Gill, and Huaming Wu, *Senior Member, IEEE*

Abstract—A wide variety of Mobile Devices (MDs) are adopted in Internet of Things (IoT) environments, resulting in a dramatic increase in the volume of task data and greenhouse gas emissions. However, due to the limited battery power and computing resources of MD, it is critical to process more data with less energy. This paper studies the Wireless Power Transfer-based Mobile Edge Computing (WPT-MEC) network system assisted by Intelligent Reflective Surface (IRS) to enhance communication performance while improving the battery life of MD. In order to maximize the Computation Energy Efficiency (CEE) of the system and reduce the carbon footprint of the MEC server, we jointly optimize the CPU frequencies of MDs and MEC server, the transmit power of Power Beacon (PB), the processing time of MEC server, the offloading time and the energy harvesting time of MDs, the local processing time and the offloading power of MD and the phase shift coefficient matrix of Intelligent Reflecting Surface (IRS). Moreover, we transform this joint optimization problem into a fractional programming problem. We then propose the Dinkelbach Iterative Algorithm with Gradient Updates (DIA-GU) to solve this problem effectively. With the help of convex optimization theory, we can obtain closed-form solutions, revealing the correlation between different variables. Compared to other algorithms, the DIA-GU algorithm not only exhibits superior performance in enhancing the system's CEE but also demonstrates significant reductions in carbon emissions.

Index Terms—Mobile Edge Computing, Intelligent Reflective Surface, Wireless Power Transfer, Energy Efficiency, Carbon Emission

1 INTRODUCTION

THE rise of the Internet of Things (IoT) has led to the deployment of a large number of MDs, including smartphones, wearables, and sensors, in various environments. For instance, smart home devices can be used to perform daily tasks, while robots deployed in factories can efficiently handle manufacturing operations. However, due to the limitations of their battery power [1], [2] and computing capacities [3], these MDs may not be able to provide adequate services in complex scenarios, such as Deep Neural Network (DNN) inference [4], [5]. In addition, how to alleviate the computing pressure on these devices and meet the stringent delay requirements of mobile users in harsh communication environments is increasingly challenging [6]. This is particularly crucial for delay-sensitive applications, e.g., intelligent manufacturing, fault diagnosis, smart supply chain, Cognitive Internet of Vehicles (CIoV) [7], big data analytics, Virtual Reality (VR) and Augmented Reality (AR). Achieving low-latency, high-bandwidth, and reliable communication is critical for these applications, and it requires innovative solutions that leverage edge computing, cloud computing, and wireless communication technologies.

Mobile Edge Computing (MEC) and Wireless Power Transfer (WPT) are emerging as promising techniques that can address the aforementioned challenges to a certain extent [8]–[10]. WPT is a non-wire contact-based technology that enables the charging of devices using clean energy sources such as solar and wind energy, and uses physical space energy carriers such as electromagnetic waves and microwaves to transmit electrical energy from the power supply side to the load side. WPT has several advantages, including the ability to charge devices anytime without plugging or unplugging, no electrical contact, support for simultaneous charging of multiple devices, and the elimination of cables, increasing the convenience and flexibility of supplying power to electrical equipment. Therefore, WPT technology can help address the issue of limited battery power of many MDs. Moreover, in the face of the dilemma of insufficient computing resources for numerous devices, MEC can effectively provide services with large bandwidth, low latency and high computing capacity. By deploying different MEC nodes, tasks that cannot be handled by a single device or cannot be processed in a timely manner are offloaded to resource-rich MEC nodes. However, the widespread deployment of MEC servers using brown energy can result in significant carbon emissions. Therefore, it is imperative to prioritize the optimization of computing resources in MEC servers within MEC scenarios to minimize the carbon emissions they generate. This field of research and practice, commonly referred to as Carbon Edge Computing or Carbon-Neutral Computing [11], aims to achieve the utmost reduction in carbon emissions from MEC servers. Moreover, there is also a need to optimize the throughput of entire systems and individual servers under energy and

- J. Du, and H. Wu are with the Center for Applied Mathematics, Tianjin University, Tianjin 300072, China. E-mail: {dujunhui_0325, whming}@tju.edu.cn.
- M. Xu is with the Shenzhen Institutes of Advanced Technology, Chinese Academy of Sciences, Shenzhen 518055, China. E-mail: liruidong@ieee.org.
- S. S. Gill is with School of Electronic Engineering and Computer Science, Queen Mary University of London, United Kingdom. Email: s.s.gill@qmul.ac.uk.

(Corresponding author: Huaming Wu)

reliability constraints [12].

In addition, it is essential to address the challenges posed by the complex communication environment in reality. In recent years, both large-scale Multiple-Input Multiple-Output (MIMO) [13] and Millimeter-Wave (MMW) communications have provided ideas for improving spectral efficiency. The former involves transmitting and receiving signals through multiple antennas at the transmitter and receiver to enhance communication quality, while the latter has a wide bandwidth but faces challenges in passing through obstacles such as buildings, and requires a large number of antennas due to the limited number of propagation paths. However, both technologies require a significant investment in terms of cost and energy consumption [14]. Therefore, there is a need to enhance communication performance with less cost and energy consumption.

Intelligent Reflective Surface (IRS) is a cutting-edge technology that can significantly improve the performance of wireless communication systems by intelligently adjusting the reflective elements on their surface. An IRS is essentially a two-dimensional super surface comprised of a large number of passive reflective elements, with each element capable of imposing a certain phase shift on the input signal. The communication environment can be improved by adjusting the angles of different reflective elements [15], thereby directing the signal to the intended receiver and mitigating interference from other sources. Compared to other communication technologies, the most significant advantage of IRS technology is that it does not require complicated signal processing operations, but simply relies on the reflection of signals to enhance its performance, resulting in lower hardware cost and computational complexity. Moreover, due to its passive nature, IRS technology is highly energy-efficient, making it a promising candidate for energy-constrained IoT applications.

In this paper, we propose the integration of WPT and IRS technologies with MEC to enhance the performance of the overall system. The objective of this approach is to minimize the carbon emissions that arise from MEC servers while simultaneously maximizing the energy efficiency of system computing in obstructed communication environments. The main **contributions** of this work are three-fold:

- *Problem Formulation:* We propose IRS technology into the WPT-MEC network to improve system performance in complex communication environments. We formulate the CEE maximization problem by jointly optimizing the CPU frequency of MD, the CPU frequency of MEC server, the processing time on MEC server, the offloading time of MD, the energy harvesting time of MD, the local processing time of MD, the offloading power of MD, the transmit power of PB and the phase shift coefficient matrix of IRS.
- *Algorithm Design:* In order to solve the objective joint fractional optimization problem, we propose the Dinkelbach Iterative Algorithm with Gradient Updates (DIA-GU). The proposed DIA-GU can achieve superior CEE performance and low system energy consumption.
- *Theoretical Analysis and Experimental Verification:* We can see that the system CEE is inversely proportional

to the CPU frequency of each MD and MEC server. Therefore, in order to improve the CEE of the system, we should appropriately reduce the CPU frequency of both. We can also get that the energy power of PB is inversely proportional to the CEE of the system, so appropriately reducing the transmit power of PB can also improve the CEE of the system. For each time block, MD needs to process tasks in the whole time block to increase the CEE of the whole system.

The remainder of this paper is structured as follows. In Section 2, we provide an overview of related work in this field, while Section 3 introduces the system model. Section 4 formulates the problem of maximizing the CEE of the system. Section 5 provides the algorithm design to obtain the optimal solutions. Simulation results are provided in Section 6. Section 7 concludes this paper and highlights future directions.

2 RELATED WORK

Due to the advantages of IRS and WPT, there have been many studies integrating them into MEC to improve the performance of their models. Table 1 identifies and compares the main elements of related works with our proposed work in terms of *Optimize computation rate*, *Reduce energy consumption*, *Enhance communication* and *Allocate computing resources*.

There are many researchers who have incorporated WPT technology into their work. Zeng *et al.* [25] introduced WPT technology and adopted a completely binary strategy to jointly optimize the mode selection (local or offload), the time allocation of energy transfer and information transfer, and the local computing speed or transmission power level to maximize the total computing rate of all users. Bi *et al.* [16] also combined the WPT technology and MEC technology. It adopts a completely binary strategy to maximize the (weighted sum) computing rate of all MD in the network by jointly optimizing single computing mode selection (local or offload) and system transmission time allocation. Huang *et al.* [18] used a completely binary strategy and deep reinforcement learning to make task offloading decisions and radio resource allocation best adapt to time-varying wireless channel conditions.

As for IRS technology, there are also many studies that introduce it into the research of MEC. For instance, Souto *et al.* [19] proposed a new method based on Particle Swarm Optimization (PSO) technology to optimize beamforming at Base Station (BS) and IRS by minimizing transmission power without Channel State Information (CSI). Yang *et al.* [20] maximized the downlink achievable rate of the user by alternately optimizing the transmit power allocation at the BS and the passive array reflection coefficient at the IRS in an iterative manner. Zhang *et al.* [21] jointly optimized the reflection coefficient and transmit covariance matrix of the IRS to maximize the capacity of a point-to-point IRS-assisted MIMO system with multiple antennas at the transmitter and receiver. Huang *et al.* [26] jointly optimized the phase-shift coefficient and the transmit power in sequential time slots to maximize the long-term energy consumption for all MDs while ensuring queue stability. Wu *et al.* [27] jointly optimized the transmit beamforming by active antenna array at

TABLE 1: The qualitative comparison of the current literature. The symbol "✓" means that the factor is taken into account, and the symbol "✗" means not taking this factor into account

Model	Optimize computation rate	Reduce energy consumption	Enhance communication	Allocate computing resources
Bi <i>et al.</i> [16]	✓	✗	✗	✓
Zhang <i>et al.</i> [17]	✓	✗	✗	✗
Huang <i>et al.</i> [18]	✓	✗	✗	✗
Souto <i>et al.</i> [19]	✗	✓	✓	✗
Yang <i>et al.</i> [20]	✓	✗	✓	✓
Zhang <i>et al.</i> [21]	✓	✗	✓	✗
Yang <i>et al.</i> [22]	✗	✓	✓	✓
Sun <i>et al.</i> [23]	✗	✓	✓	✓
Chen <i>et al.</i> [24]	✓	✗	✓	✓
This work	✓	✓	✓	✓

the AP and reflect beamforming by passive phase shifters at the IRS to minimize the total transmit power at the AP.

Unfortunately, most of the above studies assume that the computing resources of the MEC server are very large, so they all omit the computing time of data tasks on the MEC server. However, in realistic environments, MEC server does not have unlimited computing resources, and the processing time of tasks on it cannot be simply ignored. In this paper, these two points are brought into the category of model consideration for optimization. In addition, most previous literature either considers the computation bits of the system and regards the energy consumption of the system as an additional constraint, or considers the energy consumption of the system and regards the computation bits of the system as an additional constraint.

Unlike the previous works, this paper focuses on communication performance and energy consumption and proposes an IRS-assisted WPT-MEC network structure, which is more suitable for mobile scenarios. In order to evaluate the trade-off between computation bits and energy consumption, we adopt a popular performance metric called Computation Energy Efficiency (CEE) [28]–[33], which is defined as the ratio of computation bits to energy consumption for communication and computation. The uniqueness of our proposed DIA-GU method is that it not only optimizes the computation rate and reduces energy consumption, but also improves the communication efficiency when computing resource allocation is considered.

3 SYSTEM MODEL

In this section, we present an overview of the proposed system, and then provide details in the following subsections.

3.1 System Overview

As shown in Fig. 1, we consider a WPT-MEC network based on an IRS-assisted scenario, which consists of an Access Point (AP) with one MEC server, one PB, and K MDs, where each MD is equipped with a rechargeable battery and a transmittable antenna.

In this paper, we employ Time Division Multiple Access (TDMA) for each time block to divide it into several different phases. Each MD first harvests energy from the energy signal emitted by the PB, and then uses the harvested energy for task offloading and local processing. Meanwhile, we assume that the energy used by each MD for offloading and local processing in each time block does not exceed

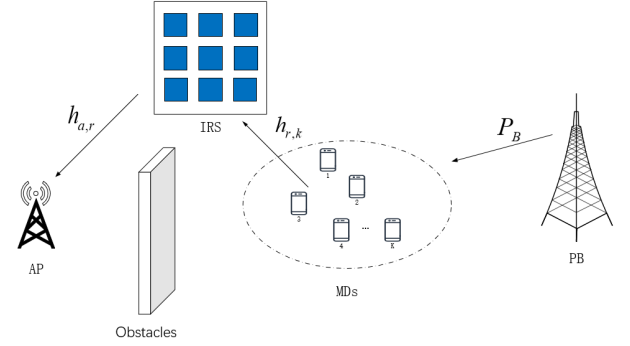


Fig. 1: An overall WPT-MEC network structure assisted by IRS.

the energy collected in each time block, so as to avoid the problem of insufficient energy in a certain time block. Many previous task offloading models are completely binary, but in this paper, we use a partial binary offloading scheme, so we assume that task data is bit-independent.

In essence, a task offloading process can be treated as a sequence of intermittent movements of task-related data across the network. As shown in Fig. 2, for any time block, there are four phases, namely, *energy harvesting*, *task offloading*, *task processing* and *download*. In the *energy harvesting* phase, each MD receives the energy signal transmitted from the PB and collects energy from it. In the *task offloading* phase, each MD offloads partial tasks to the MEC server for processing. Then, in the *task processing* phase, the task is executed on the MEC server. Finally, in the *download* phase, MEC servers typically possess robust computing capabilities, and the size of the calculation result data tends to be smaller than the task data. As a result, this paper ignores the delay involved in the MEC server transmitting the calculation results back to the MDs. During these four parts, each MD can process a part of the task data locally in each time block to improve the overall performance of the system and provide better service to the users themselves. The main notations used in this paper are summarized in Table 2.

3.2 Energy Harvesting Phase

In this phase, the PB transmits the energy signal to each MD. The energy power received by the k -th MD is as follows:

$$P_h^k = \gamma P_B h_{k,B}, \quad (1)$$

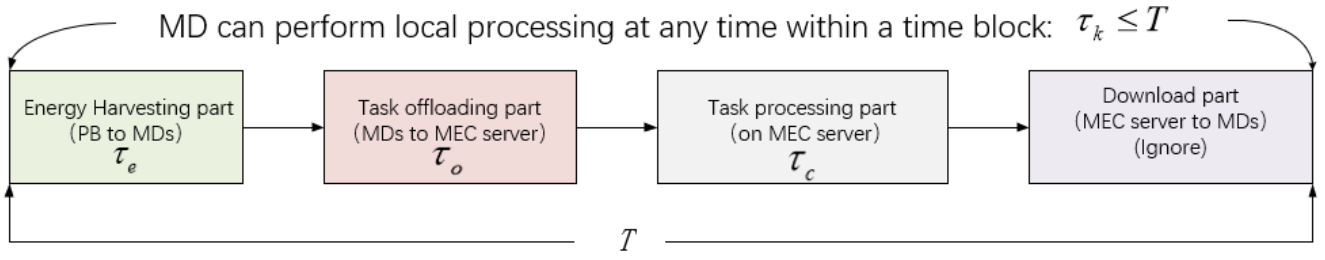


Fig. 2: Allocation of each time block.

TABLE 2: Notations and Their Definitions

Notations	Definitions
P_B	The transmit power of PB
γ	The energy conversion efficiency
P_h^k	The energy power that can be harvested of the k -th MD
θ_N	Phase shift coefficient of the N -th reflective element
$h_{r,k}$	The k -th MD to IRS link channel
$h_{a,r}$	The IRS to AP link channel
$h_{k,B}$	The k -th MD to PB link channel
θ	The phase shift coefficient matrix of IRS
h_k	The link channel from the k -th MD to the AP
j	The imaginary unit
τ_e	The energy harvesting time of MD
τ_o	The offloading time of MD
τ_c	The processing time of MEC server
τ_k	The local processing time of k -th MD
f_m	The CPU frequency on the MEC server
f_k	The CPU frequency on k -th MD processor
p_k	The offloading power of k -th MD
T	The entire time block
B	The communication bandwidth
K	The number of MDs
N	The number of IRS reflective elements
η_1	The learning rate
C_{cpu}^m	The number of CPU cycles for one bit data (MEC)
C_{cpu}^k	The number of CPU cycles for one bit data (MD)
ε_k	The ECC for the k -th MD
ε_m	The ECC for the MEC server
$f_{\max}^k$	The maximum CPU frequency of the k -th MD
$f_{\max}$	The maximum CPU frequency of the MEC server
$Q_{\min}$	The minimum amount of computational data
C_l	The carbon intensity at MEC server
C_e	The largest carbon footprint of the MEC server

where P_B is the transmit power of PB and γ is the energy conversion efficiency, $h_{k,B}$ is the channel gain between PB and k -th MD.

Therefore, the energy that can be collected from the k -th MD in the energy harvesting part is $P_h^k \tau_e$. And the energy consumed by PB at this phase is $P_B \tau_e$.

3.3 Task Offloading Phase

In this phase, since there are obstacles between MDs and AP, MDs can not communicate with AP directly or the communication efficiency between them is very low, so we can enhance the communication performance between MDs and AP by deploying IRS with N reflective elements.

The channel states between AP to IRS and MDs to IRS are modeled as quasi-stationary states, i.e., the channel states within a single time block remain unchanged, however, they may vary between different time blocks. $h_{r,k} \in C^{N \times 1}$ and $h_{a,r} \in C^{N \times 1}$ represent the k -th MD to IRS and IRS to AP link channels, respectively. Since the channels are both modeled as quasi-stationary states,

they remain invariant and can be estimated within each time block. The phase shift coefficient matrix of IRS is $\theta = \text{diag}\{e^{j\theta_1}, e^{j\theta_2}, e^{j\theta_3}, \dots, e^{j\theta_N}\}$, where the only phase shift is considered (we assume that the amplitude reflection coefficient is 1 for all reflective elements [34]). Thus, according to [35], the link channel from the k -th MD to the AP is calculated by:

$$h_k = (h_{a,r})^H \theta h_{r,k}. \quad (2)$$

Then, we employ Orthogonal Frequency Division Multiple Access (OFDM), where all MDs communicate on different orthogonal frequency bands of the same size, so the amount of communicable data between the k -th MD and the MEC server Q_k is:

$$Q_k = \tau_o B \log \left(1 + \frac{p_k |h_k|^2}{\sigma^2} \right), \quad (3)$$

where B is the channel bandwidth of the sub-band, σ^2 represents the variance of the complex Gaussian channel noise, and p_k represents the data transmission power of the k -th MD.

3.4 Task Processing Phase

Within the allocated data processing time, the amount of data that the MEC server can process at this phase is as follows:

$$Q_m = \frac{\tau_c f_m}{C_{cpu}^m}, \quad (4)$$

where f_m represents the Central Processing Unit (CPU) frequency on the MEC server, C_{cpu}^m represents the number of CPU cycles required to compute one bit of data on the MEC server.

For ease of analysis, let Q represent the number of effective computed bits on the MEC server. Since we set not only the offloading time of the MD, but also the computation time of the MEC server, Q depends not only on the total achievable data transfer amount of all the MDs, but also on the maximum amount of data that can be processed on the MEC server, that is:

$$Q = \min \left\{ Q_m, \sum_k^K Q_k \right\}. \quad (5)$$

At any time in each time block, MD can process some task data locally. The k -th MD can process the amount of data as follows:

$$Q_{loc} = \frac{\tau_k f_k}{C_{cpu}^k}, \quad (6)$$

where f_k represents the CPU frequency on the k -th MD local processor, and C_{cpu}^k represents the number of CPU cycles required to compute one bit of data at the k -th MD local processor.

The energy consumption of MEC server at this phase is obtained by the following formula [36]:

$$E_m = \varepsilon_m f_m^3 \tau_c, \quad (7)$$

where ε_m is the Energy Consumption Coefficient (ECC) of the processor chip on the MEC server, which depends on the chip structure [37].

And we also consider the carbon footprint of the MEC server, which is related to the energy consumption of the MEC server, which is expressed as $E_m C_l$ [38].

Meanwhile, the energy consumption generated by each MD during local processing is:

$$E_k = \varepsilon_k f_k^3 \tau_k, \quad (8)$$

where ε_k is the ECC of the processor chip on the k -th MD.

4 PROBLEM FORMULATION

In this section, we first formulate the joint optimization problem, and then transform it into a convex optimization problem for simplification.

4.1 Problem Definition

The CEE of the whole system is defined as the ratio between the total data throughput and the total system energy consumption.

- **Total system throughput** The total system throughput consists of two main components: i) the data processed locally by MD is $\sum_{k=1}^K \frac{\tau_k f_k}{C_{cpu}^k}$; ii) the tasks of MD is offloaded to the MEC server, and the MEC server can process part of the amount of data is $\min \left\{ \frac{\tau_c f_m}{C_{cpu}^m}, \sum_{k=1}^K \tau_o B \log_2 \left(1 + \frac{p_k |h_k|^2}{\sigma^2} \right) \right\}$;
- **Total system energy consumption** The total energy consumption of the system consists mainly of three parts: i) the energy consumption of PB is $P_B \tau_e$; ii) the energy consumption of the MEC server is $\varepsilon_m f_m^3 \tau_c$; iii) MD's local energy consumption is $\sum_{k=1}^K \varepsilon_k f_k^3 \tau_k + \sum_{k=1}^K p_k \tau_o - \sum_{k=1}^K \gamma P_B h_{k,B} \tau_e$.

Then, the system CEE can be formulated as:

$$q(f_k, f_m, \tau_e, \tau_o, \tau_c, \tau_k, p_k, P_B, \theta) = \frac{\min \left\{ \frac{\tau_c f_m}{C_{cpu}^m}, \sum_{k=1}^K \tau_o B \log_2 \left(1 + \frac{p_k |h_k|^2}{\sigma^2} \right) \right\} + \sum_{k=1}^K \frac{\tau_k f_k}{C_{cpu}^k}}{P_B \tau_e + \varepsilon_m f_m^3 \tau_c + \sum_{k=1}^K \varepsilon_k f_k^3 \tau_k + \sum_{k=1}^K p_k \tau_o - \sum_{k=1}^K \gamma P_B h_{k,B} \tau_e}. \quad (9)$$

In this paper, we jointly optimize the CPU frequency of the MD, the CPU frequency of the MEC server, the energy harvesting time of the MD, the offloading time of MD, the MEC server task processing time, the MD's local task processing time, the offloading power of MD, the transmit power of the PB, along with the scattering coefficient matrix

of the IRS to maximize the CEE of the entire system. Therefore, the optimization objective function and the problem to be optimized in this article can be expressed as $\mathcal{P}_0$:

$$\mathcal{P}_0 : \max_{\{f_k\}_{k=1}^K, \{f_m, \tau_e, \tau_o, \tau_c, \tau_k\}_{k=1}^K, \{p_k\}_{k=1}^K, P_B, \theta} q \quad (10)$$

$$\text{s.t.} \quad \min \left\{ \frac{\tau_c f_m}{C_{cpu}^m}, \sum_{k=1}^K \tau_o B \log_2 \left(1 + \frac{p_k |h_k|^2}{\sigma^2} \right) \right\} + \sum_{k=1}^K \frac{\tau_k f_k}{C_{cpu}^k} \geq Q_{\min}, \quad (10a)$$

$$p_k \tau_o + \varepsilon_k f_k^3 \tau_k \leq \gamma P_B h_{k,B} \tau_e, \quad \forall k, \quad (10b)$$

$$\tau_e + \tau_o + \tau_c \leq T, \quad (10c)$$

$$0 \leq \tau_k \leq T, \quad \forall k, \quad (10d)$$

$$0 \leq f_m \leq f_{\max}, \quad (10e)$$

$$0 \leq f_k \leq f_{\max}^k, \quad \forall k, \quad (10f)$$

$$0 \leq p_k \leq p_{\max}^k, \quad \forall k, \quad (10g)$$

$$0 \leq P_B \leq P_{\max}, \quad (10h)$$

$$\tau_e, \tau_o, \tau_c \geq 0, \quad (10i)$$

$$\varepsilon_m f_m^3 \tau_c C_l \leq C_e. \quad (10j)$$

where $Q_{\min}$ represents the minimum amount of calculation data required by all MD in the current time block, $f_{\max}$ and $f_{\max}^k$ represent the maximum CPU frequency of the MEC server and the local maximum CPU frequency of the k -th MD, $P_{\max}$ and $p_{\max}^k$ represent the maximum transmit power of PB and the local maximum offloading power of the k -th MD. Constraint (10a) is to ensure that the required amount of calculation data is completed in a time block. Constraint (10b) is to ensure that the energy consumed by the k -th MD in a time block does not exceed the harvested energy. Constraint (10c) is to ensure that the whole process of offloading to MEC server does not exceed the length of one-time block. Constraint (10d) is to ensure that the time length of local processing of the k -th MD does not exceed the length of one-time block. Constraints (10e) and (10f) are used to limit the CPU frequency of MEC server and the k -th MD, respectively, and constraints (10g) and (10h) are used to limit the offloading power of each MD and the transmit power of PB. Constraint (10i) means that each phase is non-zero in length. Constraint (10j) is to limit the carbon footprint of MEC server.

It can be seen that $\mathcal{P}_0$ is a non-convex optimization problem, which is difficult to optimize because of the coupling relationship between different parameters. Next, we will further simplify the problem.

4.2 Problem Transformation

To eliminate the coupling between MD offloading power p_k and MD data offloading length τ_o , we divide the molecular denominator of the objective function in problem $\mathcal{P}_0$ by τ_o at the same time, so the original problem $\mathcal{P}_0$ can be transformed into problem $\mathcal{P}_1$, where we let $\tau_e/\tau_o = t_e$, $\tau_c/\tau_o = t_c$, $\tau_k/\tau_o = t_k$, $1/\tau_o = t_o$.

CEE, then the following equations must hold:

$$\begin{aligned} \mathcal{P}_1 : \quad & \max_{\{f_k\}_{k=1}^K, \{f_m\}, t_e, t_o, t_c, \{t_k\}_{k=1}^K, \{p_k\}_{k=1}^K, P_B, \theta} q \\ & = \frac{\min \left\{ \frac{t_c f_m}{C_{cpu}^m}, \sum_{k=1}^K B \log_2 \left(1 + \frac{p_k |h_k|^2}{\sigma^2} \right) \right\} + \sum_{k=1}^K \frac{t_k f_k}{C_{cpu}^k}}{P_B t_e + \varepsilon_m f_m^3 t_c + \sum_{k=1}^K \varepsilon_k f_k^3 t_k + \sum_{k=1}^K p_k - \sum_{k=1}^K \gamma P_B h_{k,B} t_e} \end{aligned} \quad (11)$$

$$\begin{aligned} s.t. \quad & \min \left\{ \frac{t_c f_m}{C_{cpu}^m}, \sum_{k=1}^K B \log_2 \left(1 + \frac{p_k |h_k|^2}{\sigma^2} \right) \right\} \\ & + \sum_{k=1}^K \frac{t_k f_k}{C_{cpu}^k} \geq Q_{\min} t_o, \quad (11a) \\ & p_k + \varepsilon_k f_k^3 t_k \leq \gamma P_B h_{k,B} t_e, \forall k, \quad (11b) \\ & t_e + 1 + t_c \leq T t_o, \quad (11c) \\ & 0 \leq t_k \leq T t_o, \forall k, \quad (11d) \\ & \text{Constraints (10e) - (10h),} \\ & t_e, t_o, t_c \geq 0, \quad (11e) \\ & \varepsilon_m f_m^3 t_c C_l \leq C_e t_o. \quad (11f) \end{aligned}$$

In order to eliminate the influence of the *min* function in the objective function, we present a relaxation variable λ ($\lambda > 0$) in problem $\mathcal{P}_1$, let $\lambda = \min \left\{ \frac{t_c f_m}{C_{cpu}^m}, \sum_{k=1}^K B \log_2 \left(1 + \frac{p_k |h_k|^2}{\sigma^2} \right) \right\}$, and introduce two new constraints to get problem $\mathcal{P}_2$:

$$\begin{aligned} \mathcal{P}_2 : \quad & \max_{\{f_k\}_{k=1}^K, \{f_m\}, t_e, t_o, t_c, \{t_k\}_{k=1}^K, \{p_k\}_{k=1}^K, P_B, \theta, \lambda} q \\ & = \frac{\lambda + \sum_{k=1}^K \frac{t_k f_k}{C_{cpu}^k}}{P_B t_e + \varepsilon_m f_m^3 t_c + \sum_{k=1}^K \varepsilon_k f_k^3 t_k + \sum_{k=1}^K p_k - \sum_{k=1}^K \gamma P_B h_{k,B} t_e} \end{aligned} \quad (12)$$

$$s.t. \quad \lambda + \sum_{k=1}^K \frac{t_k f_k}{C_{cpu}^k} \geq Q_{\min} t_o, \quad (12a)$$

Constraints (11b) – (11f),

$$\frac{t_c f_m}{C_{cpu}^m} \geq \lambda, \quad (12b)$$

$$\sum_{k=1}^K B \log_2 \left(1 + \frac{p_k |h_k|^2}{\sigma^2} \right) \geq \lambda. \quad (12c)$$

It is evident that $\mathcal{P}_2$ remains a challenging nonconvex optimization problem and even worse, a difficult fractional optimization problem. To address this, we can convert $\mathcal{P}_2$ into a subtraction problem, which is prone to optimization using Dinkelbach's method [39].

Proposition 1. If $\{f_k^*\}, \{f_m^*\}, \tau_e^*, \tau_o^*, \tau_c^*, \{\tau_k^*\}, \{p_k^*\}, P_B^*, \lambda^*$ represents the optimal solution to problem $\mathcal{P}_2$, and q^* is the optimal

$$\begin{aligned} & \max_{\{f_k\}_{k=1}^K, \{f_m\}, t_e, t_o, t_c, \{t_k\}_{k=1}^K, \{p_k\}_{k=1}^K, P_B, \theta, \lambda} \lambda + \sum_{k=1}^K \frac{t_k f_k}{C_{cpu}^k} \\ & - q^* \left(P_B t_e + \varepsilon_m f_m^3 t_c + \sum_{k=1}^K \varepsilon_k f_k^3 t_k + \sum_{k=1}^K p_k - \sum_{k=1}^K \gamma P_B h_{k,B} t_e \right) \\ & = \lambda^* + \sum_{k=1}^K \frac{t_k^* f_k^*}{C_{cpu}^k} - q^* \left(P_B^* t_e^* + \varepsilon_m (f_m^*)^3 t_c^* + \sum_{k=1}^K \varepsilon_k (f_k^*)^3 t_k^* \right. \\ & \quad \left. + \sum_{k=1}^K p_k^* - \sum_{k=1}^K \gamma P_B^* h_{k,B} t_e^* \right) \end{aligned} \quad (13)$$

Proof: To provide a detailed proof of the Dinkelbach transformation, we refer to [40]. In Proposition 1, it is worth noting that the optimal parameter is identical in both $\mathcal{P}_1$ and $\mathcal{P}_2$ under the condition that the target function on the molecule exhibits concavity, while the target function in the denominator demonstrates convexity. In order to establish the concavity of the objective function in the molecular denominator, we will provide a rigorous proof by performing variable substitutions involving various parameters. ■

According to Proposition 1, the optimal solution can be obtained by using Dinkelbach iterative algorithm. The general Dinkelbach iterative algorithm determines the value of each parameter by initializing and updating the value of q , and finally ends the iteration by giving a limit on the error range. The detailed algorithmic process is summarized in Algorithm 1.

Algorithm 1 Dinkelbach iterative algorithm (DIA) for $\mathcal{P}_2$.

- 1: Initialize q and set the maximum error tolerance ε ;
 - 2: **while** true **do**
 - 3: Calculate the optimal value of $f_k, f_m, \tau_c, \tau_o, \tau_e, \tau_k, p_k, P_B$ with q value;
 - 4: Calculate a new CEE q^* ;
 - 5: **if** $|q - q^*| < \varepsilon$ **then**
 - 6: Take the obtained parameter value as the optimal value;
 - 7: **end if**
 - 8: **if** $q^* > q$ **then**
 - 9: Let $q = q^*$
 - 10: **end if**
 - 11: **end while**
-

In order to deal with the coupling relationship between different parameters, we make $x_m = t_c f_m$, $x_k = t_k f_k$ and $y_m = t_c f_m^3$, $y_k = t_k f_k^3$. In this case, $\mathcal{P}_2$ can be further transformed into $\mathcal{P}_3$:

$$\begin{aligned} \mathcal{P}_3 : \quad & \max_{\{x_k\}_{k=1}^K, x_m, t_e, t_o, \{y_k\}_{k=1}^K, y_m, \{p_k\}_{k=1}^K, P_B, \theta, \lambda} \lambda + \sum_{k=1}^K \frac{x_k}{C_{cpu}^k} - q \left(P_B t_e \right. \\ & \quad \left. + \varepsilon_m y_m + \sum_{k=1}^K \varepsilon_k y_k + \sum_{k=1}^K p_k - \sum_{k=1}^K \gamma P_B h_{k,B} t_e \right) \end{aligned} \quad (14)$$

$$s.t. \quad \lambda + \sum_{k=1}^K \frac{x_k}{C_{cpu}^k} \geq Q_{\min} t_o, \quad (14a)$$

$$p_k + \varepsilon_k y_k \leq \gamma P_B h_{k,B} \tau_e, \forall k, \quad (14b)$$

$$t_e + 1 + \sqrt{\frac{x_m^3}{y_m}} \leq T t_o, \quad (14c)$$

$$0 \leq \sqrt{\frac{x_k^3}{y_k}} \leq T t_o, \forall k, \quad (14d)$$

$$0 \leq y_m \leq (f_{\max})^2 x_m, \quad (14e)$$

$$0 \leq y_k \leq (f_{\max}^k)^2 x_k, \forall k, \quad (14f)$$

$$0 \leq y_k \leq x_k (f_{\max}^k)^2, \forall k, \quad (14g)$$

$$\text{Constraints (10h) - (10i), (11g),}$$

$$\frac{x_m}{C_{cpu}^m} \geq \lambda, \quad (14h)$$

$$\sum_{k=1}^K B \log_2 \left(1 + \frac{p_k |h_k|^2}{\sigma^2} \right) \geq \lambda, \quad (14i)$$

$$\varepsilon_m y_m C_l \leq C_e t_o. \quad (14j)$$

Theorem 1. The $\mathcal{P}_3$ is a convex optimization problem.

Proof: We can easily see that the objective function and the constraints (14a)-(14b) and (14e)-(14h) are convex functions and convex constraints. The convex properties of (14c), (14d), (14i) and (14j) are shown below.

For constraints (14c) and (14d)), let the function $H(x, y) = \sqrt{x^3/y}$, we need to prove that the function $H(x, y)$ is convex with respect to x, y , so we find the second-order partial derivative of the function with respect to x, y , and we can get the Hessian matrix as:

$$\begin{bmatrix} \frac{3}{4} \frac{1}{\sqrt{xy}} & -\frac{3}{4} \sqrt{\frac{x}{y^3}} \\ -\frac{3}{4} \sqrt{\frac{x}{y^3}} & \frac{3}{4} \sqrt{\frac{x^3}{y^5}} \end{bmatrix}.$$

We can easily get that the Hessian matrix of function $H(x, y)$ is a semi-positive definite matrix, so the function $H(x, y)$ is a convex function, so the constraints (14c) and (14d) are both convex constraints.

For (14i), what we need to prove is that $\sum_{k=1}^K B \log_2 \left(1 + \frac{p_k |h_k|^2}{\sigma^2} \right)$ is concave with respect to p_k, θ , we abstract it into a function $F(\{p_k\}_{k=1}^K, \{|h_k|\}_{k=1}^K) = \sum_{k=1}^K B \log_2 \left(1 + \frac{p_k |h_k|^2}{\sigma^2} \right)$, and first prove that F is concave with respect to $\{p_k\}_{k=1}^K, \{|h_k|\}_{k=1}^K$, and then proving that $|h_k|$ is concave with respect to θ .

According to the additivity of convex functions, to prove the convexity of function $F(\{p_k\}_{k=1}^K, \{|h_k|\}_{k=1}^K) = \sum_{k=1}^K B \log_2 \left(1 + \frac{p_k |h_k|^2}{\sigma^2} \right)$, i.e., to prove the convexity of function $F(p_k, |h_k|) = B \log_2 \left(1 + \frac{p_k |h_k|^2}{\sigma^2} \right)$, we abstract it into mathematical function $f(x, y) = \ln(1 + xy^2)$, where $x = p_k/\sigma^2$ and $y = |h_k|$. We find the second-order partial derivation of function $f(x, y) = \ln(1 + xy^2)$ about x and y , and get the Hessian matrix as:

$$\begin{bmatrix} -y^4/(1 + xy^2)^2 & 2y/(1 + xy^2)^2 \\ 2y/(1 + xy^2)^2 & (2x - 2x^2 y^2)/(1 + xy^2)^2 \end{bmatrix}.$$

What we can see is that when $xy^2 \leq 2$, the hessian matrix is a semi-negative definite matrix, i.e., $\frac{p_k |h_k|^2}{\sigma^2} \leq 2$, and since $|h_k| \leq 1$, we ignore the noise power. Thus, as long as $p_k \leq 2$ satisfies, then the hessian matrix is a semi-negative definite matrix. And in the experimental part, by adjusting the maximum power, we can see that the condition is certain, so the function $f(x, y) = \ln(1 + xy^2)$ is a concave function that satisfies.

Below we prove that the convexity of $|h_k|$ with respect to θ . According to the previous formula, we know $|h_k| = |h_{a,r}^H \theta h_{r,k}| = |h_{a,r}^H \text{diag}(e^{j\theta_1}, e^{j\theta_2}, \dots, e^{j\theta_N}) h_{r,k}|$, and for ease of proof let $h_{a,r}^H = (a_1, a_2, \dots, a_N)$, $h_{r,k} = (b_1, b_2, \dots, b_N)^H$, then we can get:

$$\begin{aligned} |h_k| &= \left| (a_1, \dots, a_N) \text{diag}(e^{j\theta_1}, e^{j\theta_2}, \dots, e^{j\theta_N}) (b_1, \dots, b_N)^H \right| \\ &= \left| (a_1, \dots, a_N) \text{diag}(\cos \theta_1, \dots, \cos \theta_N) (b_1, \dots, b_N)^H \right. \\ &\quad \left. + j (a_1, \dots, a_N) \text{diag}(\sin \theta_1, \dots, \sin \theta_N) (b_1, \dots, b_N)^H \right| \\ &= \left| (a_1 b_1 \cos \theta_1 + a_2 b_2 \cos \theta_2 + \dots + a_N b_N \cos \theta_N)^2 \right. \\ &\quad \left. + (a_1 b_1 \sin \theta_1 + a_2 b_2 \sin \theta_2 + \dots + a_N b_N \sin \theta_N)^2 \right|^{\frac{1}{2}}. \end{aligned} \quad (16)$$

To better solve for the concavity of $|h_k|$ about θ , we set the function $h(\theta)$ and prove its concavity as follows:

$$\begin{aligned} h(\theta) &= (a_1 b_1 \cos \theta_1 + a_2 b_2 \cos \theta_2 + \dots + a_N b_N \cos \theta_N)^2 \\ &\quad + (a_1 b_1 \sin \theta_1 + a_2 b_2 \sin \theta_2 + \dots + a_N b_N \sin \theta_N)^2 \\ &= \sum_{i \neq j} (2a_i b_i \cos \theta_i a_j b_j \cos \theta_j + 2a_i b_i \sin \theta_i a_j b_j \sin \theta_j) + C \\ &= \sum_{i \neq j} (2a_i b_i a_j b_j \cos(\theta_i - \theta_j)) + C, \end{aligned} \quad (17)$$

where C is a constant, let function $h_{ij}(\theta) = \cos(\theta_i - \theta_j)$, get the second-order derivative of function $h_{ij}(\theta)$ about θ_i and θ_j , and get its Hessian matrix H as follows:

$$\begin{bmatrix} 0 & & & & \\ & \ddots & & & \\ & & -\cos(\theta_i - \theta_j) & \cdots & \cos(\theta_i - \theta_j) \\ & & \vdots & \ddots & \vdots \\ & & \cos(\theta_i - \theta_j) & \cdots & -\cos(\theta_i - \theta_j) \\ & & & & \ddots & \\ & & & & & 0 \end{bmatrix}$$

From the definition of a semi-negative definite matrix, we can easily get that for any non-zero vector x , there is $x^T H x \leq 0$. We can obtain that the matrix is a semi-negative definite matrix. So the function $h_{ij}(\theta)$ is a concave function, according to the additivity of the concave function, we can get that the function $h(\theta)$ is a concave function, and because for the function $x^{\frac{1}{2}}$, we can easily get this function about x being a concave function, So we can get that $|h_k|$ about θ is concave.

In summary, this theorem is proved. ■

5 PROPOSED SOLUTIONS

5.1 Problem Solving

Theorem 2. Given non-negative Lagrangian factor $\alpha, \beta, M, H, N_1, N_2, N_3$ we can obtain the expression of the optimal value of

some parameters by solving the following Lagrangian function L as shown in Eq. (15).

By making the partial derivative of each parameter equal to zero, we can derive the optimal solution for some of the parameter variables, as follows [41]:

$$f_k^* = \left[\frac{3\beta_k}{2\left(\frac{\alpha_0+1}{C_{cpu}^k} + M_k(f_{\max}^k)^2\right)} \right]^+ = \left[\left(\frac{\beta_k}{2(q\varepsilon_k + \alpha_k\varepsilon_k + M_k)} \right)^{\frac{1}{3}} \right]^+ \quad (18)$$

$$f_m^* = \left[\frac{3\beta_0}{2\left(M_0f_{\max}^2 + \frac{N_1}{C_{cpu}^m}\right)} \right]^+ = \left[\left(\frac{\beta_0}{2(q\varepsilon_m + M_0 + N_3\varepsilon_m C_l)} \right)^{\frac{1}{3}} \right]^+ \quad (19)$$

where $[x]^+ = \max\{x, 0\}$. For the purpose of reducing the number of parameters, according to the above two formulas, we can obtain:

$$\beta_0 = \sqrt{\frac{4\left(M_0f_{\max}^2 + \frac{N_1}{C_{cpu}^m}\right)^3}{27(q\varepsilon_m + M_0 + N_3\varepsilon_m C_l)}}, \quad (20)$$

$$\beta_k = \sqrt{\frac{4\left(\frac{\alpha_0+1}{C_{cpu}^k} + M_k(f_{\max}^k)^2\right)^3}{27(q\varepsilon_k + \alpha_k\varepsilon_k + M_k)}}, \quad (21)$$

Then by partial derivation of the remaining parameters, we can also get:

$$t_e^* = \frac{H_0}{q \sum_{k=1}^K \gamma h_{k,B} + \sum_{k=1}^K \alpha_k \gamma h_{k,B} - q} = h, \quad (22)$$

$$P_B^* = \left[\frac{\beta_0}{q \sum_{k=1}^K \gamma h_{k,B} + \sum_{k=1}^K \alpha_k \gamma h_{k,B} - q} \right]^+, \quad (23)$$

$$p_k^* = \left[\frac{N_2 B}{\ln 2(q + \alpha_k + H_k)} - \frac{\sigma^2}{h_k^2} \right]^+, \quad (24)$$

And we can also get: $1 + \alpha_0 = N_1 + N_2$.

Theorem 3. In fact, in order to get a larger system CEE, the two terms represented by

$$\lambda = \min \left\{ \frac{t_c f_m}{C_{cpu}^m}, \sum_{k=1}^K B \log_2 \left(1 + \frac{p_k |h_k|^2}{\sigma^2} \right) \right\}$$

should be equal, for all the tasks offloaded to the MEC server, the MEC server can process them.

Proof: Suppose $\{f_k^*\}, \{f_m^*\}, t_e^*, t_o^*, t_c^*, \{t_k^*\}, \{P_k^*\}, q^*, |h_k^*|$ are the optimal solution to the original problem, because the two terms represented by $\lambda = \min \left\{ \frac{t_c f_m}{C_{cpu}^m}, \sum_{k=1}^K B \log_2 \left(1 + \frac{p_k |h_k|^2}{\sigma^2} \right) \right\}$ are not equal, so we may as well make the former term greater than the latter term, that is $\frac{t_c f_m}{C_{cpu}^m} > \sum_{k=1}^K B \log_2 \left(1 + \frac{p_k |h_k|^2}{\sigma^2} \right)$.

Since q^* represents the optimal solution of the system CEE, we can get $\frac{t_c f_m}{C_{cpu}^m} > \sum_{k=1}^K B \log_2 \left(1 + \frac{p_k^* |h_k^*|^2}{\sigma^2} \right)$. In contrast, we can set up another set of solutions $f_k^* = f_k, t_e^* = t_e, t_o^* = t_o, t_c^* = t_c, t_k^* = t_k^*, P_k^* = P_k, |h_k^*| = |h_k|$, and we set $\frac{t_m f_m}{C_{cpu}^m} = \sum_{k=1}^K B \log_2 \left(1 + \frac{P_k^* |h_k^*|^2}{\sigma^2} \right)$, then we can get $\frac{t_m f_m}{C_{cpu}^m} = \sum_{k=1}^K B \log_2 \left(1 + \frac{P_k^* |h_k^*|^2}{\sigma^2} \right) = \sum_{k=1}^K B \log_2 \left(1 + \frac{P_k^* |h_k^*|^2}{\sigma^2} \right) < \frac{t_m f_m}{C_{cpu}^m}$, then we can also get $f_m^* < f_m$.

From Eq. (19), we know that the value of q is inversely proportional to f_m , so the new set of solutions q^* is greater than the optimal solution q^* , resulting in a contradiction. Therefore, we can draw the conclusion that q is large when the two terms of $\lambda = \min \left\{ \frac{t_c f_m}{C_{cpu}^m}, \sum_{k=1}^K B \log_2 \left(1 + \frac{p_k |h_k|^2}{\sigma^2} \right) \right\}$ are equal. ■

By proving Theorem 3, we can derive the optimal solution for the MEC server processing time τ_c :

$$\tau_c^* = \frac{\tau_o^* B \log_2 \left(1 + \frac{P_k^* |h_k^*|^2}{\sigma^2} \right)}{f_m^*}. \quad (25)$$

Secondly, when β is greater than zero, according to the Slater condition, we can obtain the following conclusion:

$$\begin{cases} \tau_k^* = T \\ \tau_e^* + \tau_o^* + \tau_c^* = T \end{cases} \quad (26)$$

From the above equation and Theorem 3, we can get:

$$\tau_c^* = T - (h + 1) \tau_o^*, \quad (27)$$

$$\frac{\tau_c^* f_m^*}{C_{cpu}^m} = \sum_{k=1}^K \tau_o^* B \log_2 \left(1 + \frac{p_k^* |h_k^*|^2}{\sigma^2} \right), \quad (28)$$

$$\begin{aligned} L = & \lambda + \sum_{k=1}^K \frac{x_k}{C_{cpu}^k} - q \left(P_B t_e + \varepsilon_m y_m + \sum_{k=1}^K \varepsilon_k y_k + \sum_{k=1}^K p_k - \sum_{k=1}^K \gamma P_B h_{k,B} t_e \right) + \alpha_0 \left(\lambda + \sum_{k=1}^K \frac{x_k}{C_{cpu}^k} - Q_{\min} t_o \right) \\ & + \sum_{k=1}^K \alpha_k (\gamma P_B h_{k,B} t_e - p_k - \varepsilon_k y_k) + \beta_0 \left(T t_o - t_e - 1 - \sqrt{\frac{x_m^3}{y_m}} \right) + \sum_{k=1}^K \beta_k \left(T t_o - \sqrt{\frac{x_k^3}{y_k}} \right) + M_0 ((f_{\max}^k)^2 x_m - y_m) \\ & + \sum_{k=1}^K M_k ((f_{\max}^k)^2 x_k - y_k) + \sum_{k=1}^K H_k (p_{\max}^k - p_k) + H_0 (P_{\max} - P_B) + N_1 \left(\frac{x_m}{C_{cpu}^m} - \lambda \right) \\ & + N_2 \left(\sum_{k=1}^K B \log_2 \left(1 + \frac{p_k |h_k|^2}{\sigma^2} \right) - \lambda \right) + N_3 (C_e t_o - \varepsilon_m y_m C_l) \end{aligned} \quad (15)$$

So we can get the following solutions:

$$\tau_o^* = \frac{\frac{T f_m^*}{C_{cpu}^m}}{\frac{(h+1)f_m^*}{C_{cpu}^m} + \sum_{k=1}^K B \log_2 \left(1 + \frac{p_k^* |h_k^*|^2}{\sigma^2} \right)}, \quad (29)$$

$$\tau_e^* = T - \tau_o^* - \tau_c^*. \quad (30)$$

Remark: Through Eqs. (18) and (19), we can see that the system CEE is inversely proportional to the CPU frequency of each MD and MEC server. Therefore, in order to improve the CEE of the system, we should appropriately reduce the CPU frequency of both. As can be seen from Eq. (23), we can also get that the energy power of PB is inversely proportional to the CEE of the system, so appropriately reducing the transmit power of PB can also improve the CEE of the system. As can be seen from Eq. (24), only when the state of the channel is better, MD will choose to carry out data offloading. As can be seen from Eq. (26), for each time block, MD needs to process tasks in the whole time block to increase the CEE of the whole system. Without the loss of generality, we assume that each MD can perform local execution in the whole time block [16], [18], in the numerical experiment part, we will also prove the rationality of the formula through experiments.

5.2 DIA-GU Algorithm

In this paper, we propose a novel DIA-GU algorithm that builds upon the Dinkelbach general iterative algorithm, incorporating the concept of alternating iterations through the utilization of the Lagrange multiplier method. In addition to updating the q value with the optimal value of each parameter variable, we also iteratively update the Lagrangian factor within each cycle to enhance the system's overall performance. By integrating these techniques, we attain the optimal value of the objective optimization function.

First and foremost, it is important to note that the reflection coefficient vector of the IRS serves as a parameter solely linked to the data throughput of the system. Therefore, our optimization efforts can focus exclusively on enhancing the data throughput by optimizing the reflection coefficient vector of the IRS. Moreover, according to Eq. (9), we can deduce that the data throughput of the system is directly proportional to the magnitude of the channel coefficient $|h_k|$ when other parameters are held constant. Furthermore, Eq. (17) indicates that to maximize $|h_k|$, it is optimal to set the reflection angles of all the IRS elements to be equal. This configuration enables the attainment of the maximum achievable data throughput.

It is essential to note that these conclusions are applicable specifically to the scenario described in this paper, wherein the channel states between the MD and IRS, as well as between the IRS and MEC server, are known in each time slot. In such a case, obtaining the optimal IRS reflection coefficient vector is relatively straightforward. Similarly, this conclusion can be extended to cases involving a direct link between the MD and MEC servers, where the attestation process follows a similar structure to that of Eq. (17).

The detailed process of the proposed algorithm is as described in Algorithm 2. In each iteration, we can calculate the optimal values of $f_k, f_m, \tau_c, \tau_o, \tau_e, \tau_k, p_k, P_B$ with

q value, then we fix these parameters and do a gradient update on Lagrangian factor, and then we can obtain new q^* value from these optimal values. We continue the above steps with q^* instead of q value until the algorithm converges.

Algorithm 2 Dinkelbach iterative algorithm with Gradient updates (DIA-GU) for $\mathcal{P}_3$

- 1: Initialize q, θ ($\theta_1 = \theta_2 = \dots = \theta_N$), Lagrangian factors and learning rate η_1 , and set the maximum error tolerance ε ;
- 2: **while** true **do**
- 3: Fixed the Lagrange factor, and calculate the optimal value of $f_k, f_m, \tau_c, \tau_o, \tau_e, \tau_k, p_k, P_B$ with q value;
- 4: Based on the new parameter values already obtained in the above steps, we can calculate a new CEE q^* ;
- 5: **if** $q^* > q$ **then**
- 6: Let $q = q^*$;
- 7: **end if**
- 8: Keep q unchanged, only update the Lagrange factor:
- 9: $\alpha = \alpha - \eta_1 * \frac{\partial L}{\partial \alpha}, M = M - \eta_1 * \frac{\partial L}{\partial M}, M_0 = M_0 - \eta_1 * \frac{\partial L}{\partial M_0}, \beta = \beta - \eta_1 * \frac{\partial L}{\partial \beta}$;
- 10: $N_1 = N_1 - \eta_1 * \frac{\partial L}{\partial N_1}, N_2 = N_2 - \eta_1 * \frac{\partial L}{\partial N_2}, H = H - \eta_1 * \frac{\partial L}{\partial H}, N_3 = N_3 - \eta_1 * \frac{\partial L}{\partial N_3}$;
- 11: **if** $|q - q^*| < \varepsilon$ **then**
- 12: Take the obtained parameter value as the optimal value;
- 13: **end if**
- 14: **end while**

5.2.1 Complexity Analysis

The algorithm proposed in this paper is comprised of a loop iteration. Specifically, the variable I_1 represents the number of loops. Therefore, the computational complexity of the proposed algorithm can be expressed as $O(I_1)$.

5.2.2 Convergence Analysis

Let $Q(A, B)$ denote the value of the target function, where $A = \{f_k, f_m, \tau_c, \tau_o, \tau_e, \tau_k, p_k, P_B\}$, and $B = \{\alpha, \beta, M, H, N_1, N_2, N_3\}$. In the i -th iteration, we can obtain:

$$Q(A^i, B^i) \stackrel{(a)}{\leq} Q(A^{i+1}, B^i) \stackrel{(b)}{\leq} Q(A^{i+1}, B^{i+1}). \quad (31)$$

where inequality (a) holds true as the proposed algorithm obtains the optimal value for each parameter in every iteration, ensuring that it does not decrease throughout the iterative process. Inequality (b) arises from the utilization of gradient updates to modify each Lagrangian factor, thereby bringing the objective function's value closer to its optimal state. Consequently, the proposed algorithm gradually approaches the optimal value of the objective function and achieves convergence during the iterative process.

6 PERFORMANCE EVALUATION

In this section, we demonstrate the superiority of the proposed algorithm by conducting a large number of single-machine simulation experiments through Python.

6.1 Parameter Settings

Unless otherwise specified, the basic simulation parameters are given as shown in Table 3, and for $h_{r,k} \in \mathbb{C}^{N \times 1}$, we set it to 0.2. For $h_{a,r} \in \mathbb{C}^{N \times 1}$ and $h_{k,B} \in \mathbb{C}^{K \times 1}$, we set them to 0.2.

TABLE 3: Parameter Settings

Parameter	Description	Value
T	The entire time block	1 second
B	The communication bandwidth	1 MHz
ε	The maximum error tolerance	0.01
P_{max}	The PB's maximum transmit power	5.0 W
p_{max}^k	The k -th MD's maximum offload power	0.02 W
γ	The energy conversion efficiency	0.1
K	No. of MDs	4
N	No. of IRS reflective elements	12
η_1	Learning rate	0.001
ε_k	The ECC of the k -th MD	10^{-26}
ε_m	The ECC of MEC server	10^{-28}
f_{max}^k	Maximum CPU frequency of k -th MD	5×10^8 Hz
f_{max}	Maximum CPU frequency of MEC server	10^{10} Hz
Q_{min}	The minimum computation bits	5×10^5 bit
C_{cpu}^m, C_{cpu}^k	No. of CPU cycles for one bit data	1000 cycles/bit
C_l	The carbon intensity at MEC server	300 g/kWh
C_e	The largest carbon footprint of the MEC server	7×10^{-3} g

6.2 Baselines

We compare the other iterative algorithms in system CEE, mainly by comparing the following four algorithms:

- *Dinkelbach Iterative Algorithm with Gradient Updates (DIA-GU)*: The iterative algorithm proposed in this paper leverages the characteristics of both the Dinkelbach algorithm and the Lagrange multiplier method.
- *0.5*T local computing*: The local processing time of the algorithm is 0.5 time slots.
- *MAX computing*: The maximum calculation frequency of the MD is directly obtained in each time slot.
- *to random computing*: The duration of offloading time in each time slot is determined randomly.
- *tc random computing*: The processing time of the MEC server in each time slot is randomly determined.

6.3 Experimental Analysis

As shown in Fig. 3, we conduct a comparison between the system's data throughput under equal IRS reflectance coefficients and unequal IRS reflectance coefficients. The experimental results demonstrate that utilizing equal IRS reflection coefficients leads to higher system data throughput. This outcome serves as empirical evidence that corroborates the validity of the aforementioned theoretical proof. Furthermore, we perform experiments with different angles while employing equal IRS reflection coefficients. Notably, the results reveal that, in the presence of equal IRS reflection coefficients, the data throughput remains consistent across varying angles. This finding further affirms the accuracy of the earlier theoretical proof.

As shown in Fig. 4, by comparing the trend of system CEE under different f_{max}^k , we can see that as the maximum computational frequency of the local processor increases, the system CEE increases. By Eq. (18), we can implicitly

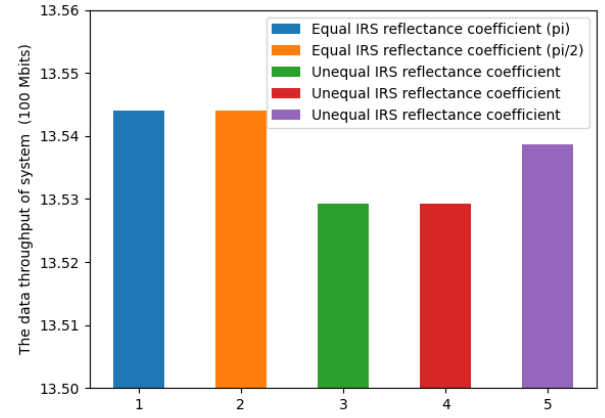


Fig. 3: System data throughput under different IRS reflectance coefficient

get this conclusion. First, we get from the closed solution of the optimal calculation frequency of the local processor: the denominator in the optimal solution of the parameter variable is related to f_{max}^k , note that the f_{max}^k in the denominator is inversely proportional to the optimal calculation frequency of the local processor, and in the second equation of this equation we can get that the q value is also inversely proportional to the optimal calculation frequency of the local processor. Therefore, we can simply conclude that for the increase of f_{max}^k , since the optimal solution for f_k is non-additive, the consequent increase in the CEE of the system is the same as our simulation results. Therefore, in the deployment of real scenarios, under the same conditions, we can appropriately increase the maximum computing frequency of the local processor to obtain better system performance.

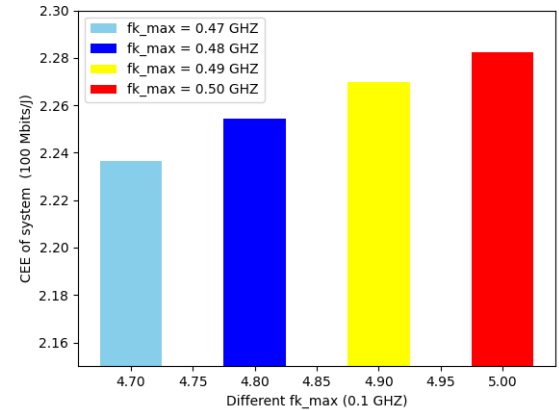


Fig. 4: System CEE under different f_{max}^k

As shown in Fig. 5, we study the trend of system CEE under different numbers of MDs. The results demonstrate a positive correlation, indicating that as the number of MDs increases, the CEE of the entire system also increases. This observation aligns with our intuitive understanding and underscores the scalability of the scenario investigated in this paper.

By analyzing the trend of the system's CEE across different P_{max} , as depicted in Fig. 6, we can draw the conclusion

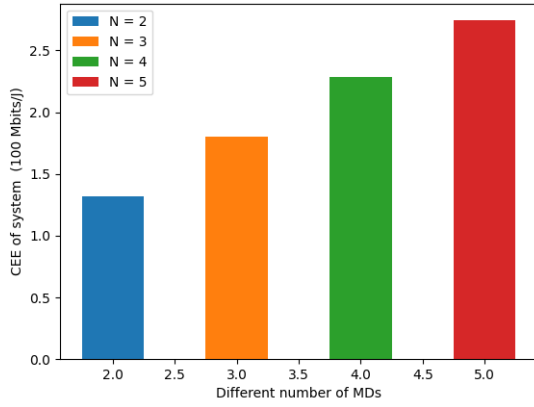


Fig. 5: System CEE under different number of MDs

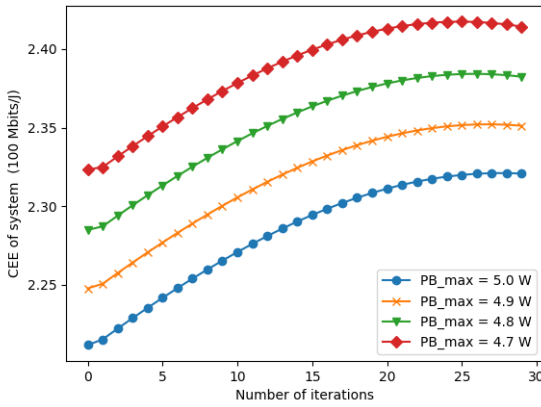


Fig. 6: System CEE under different P_{max}

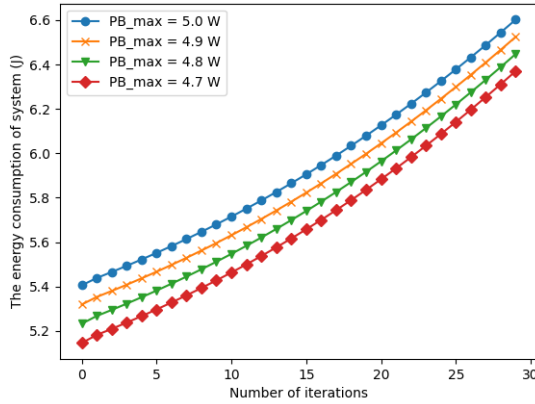


Fig. 7: System energy consumption under different P_{max}

that as the maximum transmit power of PB increases, the system CEE decreases. This observation can be explained both theoretically and experimentally.

- From a theoretical standpoint, we can attribute this phenomenon to the non-subtractive nature of the optimal transmit power of the PB. According to Eq. (23), an inverse relationship exists between the q value and the optimal transmission power of the PB. Therefore, as P_{max} increases, the q value decreases.
- From an experimental perspective, we observe that

the optimal transmission power of the PB is non-subtractive, leading to an increase in PB's energy consumption. Consequently, the total energy consumption of the system also rises, as evidenced in Fig. 7. With the increase in P_{max} , a clear upward trend in total energy consumption is evident.

For data throughput, as shown in Fig. 8, we can see that with the increase of P_{max} , the data throughput is decreasing accordingly. Combining the results of the aforementioned experiments, where the numerator (data throughput) decreases and the denominator (energy consumption) increases, we can infer a decline in system performance, specifically a decrease in the q value.

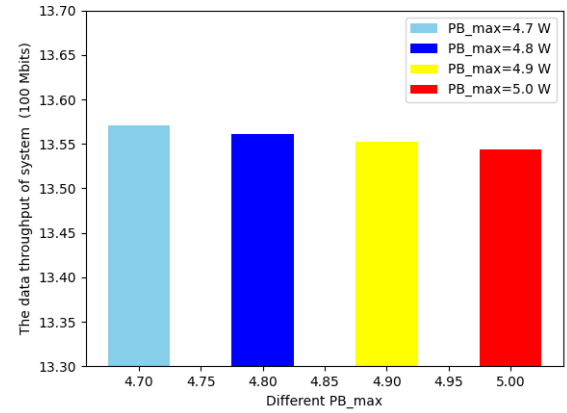


Fig. 8: System data throughput under different P_{max}

As shown in Fig. 9, we investigated the correlation between $\varepsilon_m/\varepsilon_k$ (m/k) and the system's CEE. We can conclude that as $\varepsilon_m/\varepsilon_k$ increases, the performance of the entire system declines. In the case of ECC, this ratio can be perceived as an indicator of the energy efficiency of the processor itself. Alternatively, we can examine this phenomenon from a different perspective. As the ratio increases, if the ECC of the MEC server remains the same, it implies a reduction in the ECC of the MEC server. In other words, when the ECC of the MD decreases, it indicates that the MD consumes less energy under the same conditions. To illustrate this point further, let's consider an extreme scenario where the ratio equals 1, signifying an equal energy performance between the two entities. In this case, for a given task and calculation frequency, the energy consumed by both the MD and the MEC server is the same. However, opting to offload the task to the MEC server incurs additional energy consumption due to the data transmission process. Consequently, the system gradually leans toward local task processing to conserve energy. However, we are aware that local processing capacity is limited, leading to a decline in the overall system performance. Through this experiment, we can infer that $\varepsilon_m/\varepsilon_k$ effectively represents the energy performance of both the MEC server and the MD device to a certain extent. When the performance gap between the two entities becomes too narrow, it results in a deterioration of the overall system performance.

We also study the trend of system CEE under different EH factor γ , as shown in Fig. 10, we can clearly see that with the growth of EH factor, system CEE also grows, we can get

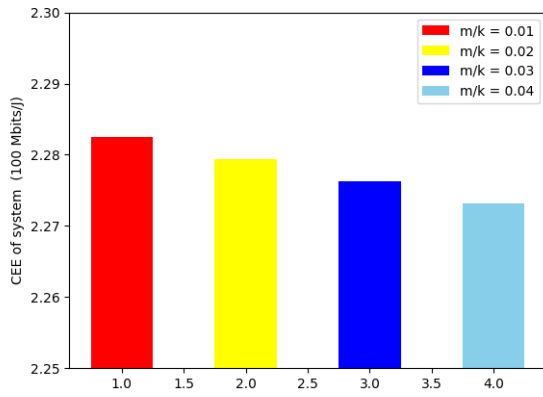


Fig. 9: System CEE under different $\varepsilon_m/\varepsilon_k$ (m/k)

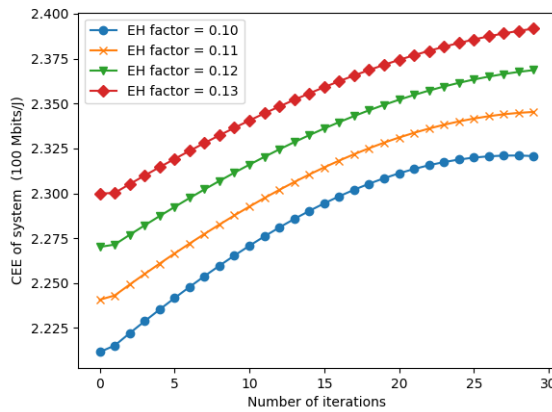


Fig. 10: System CEE under different EH factor γ

this conclusion through the expression of system CEE. The EH factor affects the energy that can be obtained by MD in each time slot, and under other conditions being equal, as the energy obtained in each time slot increases, the energy consumed by the entire system is decreasing, and finally, the CEE of the system is increasing.

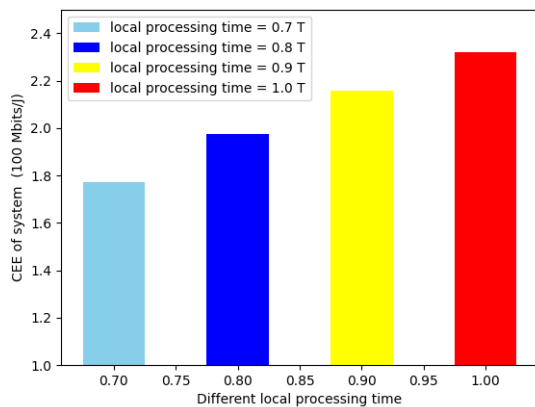


Fig. 11: System CEE under different proportions of local processing time

As shown in Fig. 11, we study the change of system CEE under different proportions of local processing time, and we see that as the local processing time increases, the

system CEE also increases, and the experiment proves that we directly use the entire time slot as the local optimal processing time.

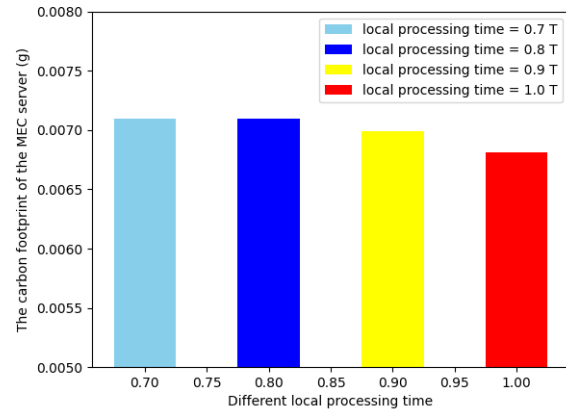


Fig. 12: Carbon emissions from MEC server under different proportions of local processing time

As shown in Fig. 12, we not only compare the system performance changes caused by different local processing durations, but also analyze the carbon emissions caused by MEC servers under different local processing durations. The result shows that using the entire time slot as the local processing time can effectively reduce the carbon emission of the MEC server, because the increase of local processing time makes the MEC server need to process fewer tasks, which naturally reduces carbon emissions, so from the perspective of reducing carbon emissions, the validity of the previous theory can also be proved.

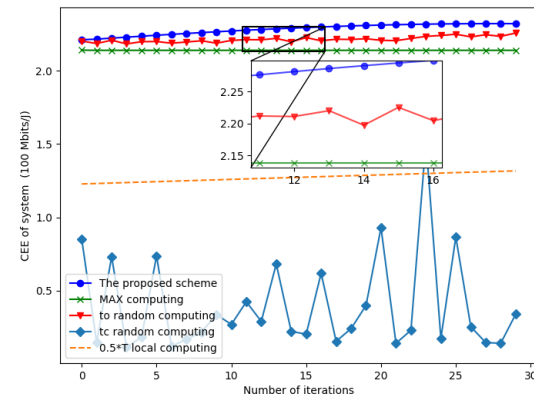


Fig. 13: System CEE under different algorithms

Fig. 13 illustrates the comparison between our DAI-GU algorithm and four baseline algorithms. The scenario described in this paper necessitates a minimum amount of processed data per time slot, making local processing alone inadequate. Consequently, the baseline algorithms do not include a complete local processing algorithm for comparison. However, we introduce a local algorithm that serves as the benchmark, with a local processing time of only 0.5 time slots. Additionally, we compare the MEC server processing time random algorithm, the task offloading time random algorithm, and an algorithm that utilizes maximum

computing resources. To provide a comprehensive analysis, we conduct detailed experiments separately. Through these experiments, we observe the superior performance of the proposed algorithm in enhancing the system's CEE.

As depicted in Fig. 14, we also compare the system throughput among different algorithms. Through simulation experiments, it is evident that the proposed algorithm achieves the maximum data throughput, on par with the *MAX computing* algorithm. Furthermore, in comparison to the *MAX computing* algorithm, the proposed algorithm demonstrates the capability to achieve approximately 80% of its maximum data throughput.

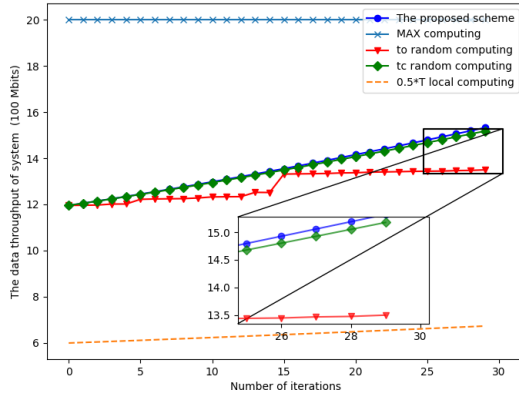


Fig. 14: System data throughput under different algorithms

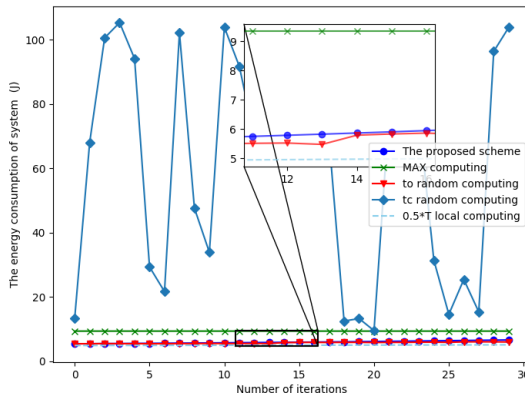


Fig. 15: System energy consumption under different algorithms

In Fig. 15, we conducted a comparative analysis of the system's energy consumption. Initially, we observed that the energy consumption exhibited significant oscillations and was excessively high due to the *tc random computing* algorithm. The results revealed that the *0.5*T local computing* algorithm yielded the lowest system energy consumption. Comparatively, this algorithm only increased energy consumption by a modest 20% when compared to the algorithm proposed in this paper. Furthermore, when compared to the *MAX computing* algorithm, the proposed algorithm reduced energy consumption by an impressive 35%. Consequently, it can be concluded that the algorithm proposed in this paper effectively combines the objectives of enhancing data

throughput and reducing system energy consumption. This amalgamation results in improved system performance, offering a compelling solution for optimizing energy consumption while maintaining or enhancing data throughput.

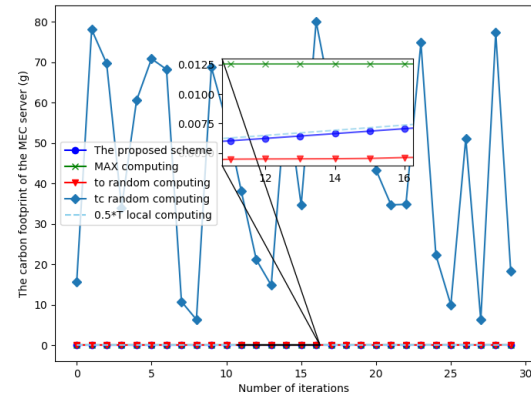


Fig. 16: Carbon emissions of MEC server under different algorithms

As shown in Fig. 16, we present a comparison of the carbon emissions of MEC servers across various algorithms. In contrast to the *MAX computing* algorithm, the proposed algorithm exhibits a noteworthy reduction in carbon emissions of approximately 30%. On the other hand, when compared to the *to random computing* algorithm, which demonstrates the lowest carbon emissions, the proposed algorithm only marginally increases carbon emissions by less than 40%. However, it is important to note that this increase in carbon emissions is accompanied by a significant improvement in system performance.

7 CONCLUSIONS AND FUTURE WORK

In this paper, we studied the maximization of CEE of the WPT-MEC network assisted by IRS. We jointly optimized the CPU frequency of MD, the CPU frequency of the MEC server, the transmit power of PB, the processing time on the MEC server, the offloading time of MD, the energy harvesting time of MD, the local processing time of MD, the offloading power of MD and the phase shift coefficient matrix of IRS. In order to solve the problem of joint optimization fraction, we proposed an iterative algorithm based on Dinkelbach theory and improved the algorithm to make it more suitable for the application scenario. In order to further improve the performance of the system, we proposed the DIA-GU algorithm. Compared with other algorithms, the DIA-GU algorithm can not only perform better in improving the CEE of the system, but also in reducing carbon emissions. Moreover, we can get many beneficial insights from the closed-form solution of each parameter. For instance, the system CEE increases as the MD local processor and the CPU frequency of the MEC server decrease, and the total amount of data offloaded from all MDs should be equal to the maximum amount of data that the MEC server can process during the MEC server processing phase, and each MD should use the maximum allowable time to process the local task data.

In future work, it would be beneficial to incorporate the consumption of stored energy per time slot to enhance the performance of the model. For the communication model, the OFDM communication model is selected, and in the future, communication models such as NOMA that save spectrum resources can also be considered. In addition, it is worth emphasizing that the relationship between the two distinct channel states and the IRS reflection angle is interconnected. The current scenario is based on a quasi-stationary channel model for solving, and future research will explore more intricate scenarios, e.g., the correlation model between the two channel states and the IRS reflection angle.

ACKNOWLEDGMENTS

This work is supported by the National Natural Science Foundation of China (No. 62071327, 62102408), Tianjin Science and Technology Planning Project (No. 22ZYYJC00020), Shenzhen Science and Technology Program (No. RCBS20210609104609044), Chinese Academy of Sciences President's International Fellowship Initiative (Grant No. 2023VTC0006) and Shenzhen Industrial Application Projects of undertaking the National key R & D Program of China (No. CJGJZD20210408091600002).

REFERENCES

- [1] Z. Xu, J. Zhou, W. Zheng, Y. Wang, and M. Guo, "Exploiting big, little batteries for software defined management on mobile devices," *IEEE Transactions on Mobile Computing*, vol. 21, no. 6, pp. 1998–2012, 2022.
- [2] J. Zhou, Z. Xu, W. Zheng, and Y. Wang, "Capman: Cooling and active power management in big, little battery supported devices," in *2020 IEEE 40th International Conference on Distributed Computing Systems (ICDCS)*, 2020, pp. 1009–1019.
- [3] K. W. Choi, L. Ginting, A. A. Aziz, D. Setiawan, J. H. Park, S. I. Hwang, D. S. Kang, M. Y. Chung, and D. I. Kim, "Toward realization of long-range wireless-powered sensor networks," *IEEE Wireless Communications*, vol. 26, no. 4, pp. 184–192, 2019.
- [4] M. Xue, H. Wu, R. Li, M. Xu, and P. Jiao, "Eosdnn: An efficient offloading scheme for dnn inference acceleration in local-edge-cloud collaborative environments," *IEEE Transactions on Green Communications and Networking*, vol. 6, no. 1, pp. 248–264, 2022.
- [5] S. S. Gill, M. Xu, C. Ottaviani, P. Patros, R. Bahsoon, A. Shaghghi, M. Golec, V. Stankovski, H. Wu, A. Abraham, M. Singh, H. Mehta, S. K. Ghosh, T. Baker, A. K. Parlikad, H. Lutfiyya, S. S. Kanhere, R. Sakellariou, S. Dustdar, O. Rana, I. Brandic, and S. Uhlig, "AI for next generation computing: Emerging trends and future directions," *Internet of Things*, vol. 19, p. 100514, aug 2022.
- [6] G. Qu, N. Cui, H. Wu, R. Li, and Y. Ding, "Chainfl: A simulation platform for joint federated learning and blockchain in edge/cloud computing environments," *IEEE Transactions on Industrial Informatics*, vol. 18, no. 5, pp. 3572–3581, 2022.
- [7] X. Xu, B. Shen, X. Yin, M. R. Khosravi, H. Wu, L. Qi, and S. Wan, "Edge server quantification and placement for offloading social media services in industrial cognitive iot," *IEEE Transactions on Industrial Informatics*, vol. 17, no. 4, pp. 2910–2918, 2021.
- [8] P. Ramezani and A. Jamalipour, "Toward the evolution of wireless powered communication networks for the future internet of things," *IEEE Network*, vol. 31, no. 6, pp. 62–69, 2017.
- [9] K. W. Choi, L. Ginting, A. A. Aziz, D. Setiawan, J. H. Park, S. I. Hwang, D. S. Kang, M. Y. Chung, and D. I. Kim, "Toward realization of long-range wireless-powered sensor networks," *IEEE Wireless Communications*, vol. 26, no. 4, pp. 184–192, 2019.
- [10] S. Bi, C. K. Ho, and R. Zhang, "Wireless powered communication: opportunities and challenges," *IEEE Communications Magazine*, vol. 53, no. 4, pp. 117–125, 2015.
- [11] L. Gu, W. Zhang, Z. Wang, D. Zeng, and H. Jin, "Service management and energy scheduling toward low-carbon edge computing," *IEEE Transactions on Sustainable Computing*, vol. 8, no. 1, pp. 109–119, 2023.
- [12] J. Zhou, K. Cao, X. Zhou, M. Chen, T. Wei, and S. Hu, "Throughput-conscious energy allocation and reliability-aware task assignment for renewable powered in-situ server systems," *IEEE Transactions on Computer-Aided Design of Integrated Circuits and Systems*, vol. 41, no. 3, pp. 516–529, 2022.
- [13] S. Hu, F. Rusek, and O. Edfors, "Beyond massive mimo: The potential of data transmission with large intelligent surfaces," *IEEE Transactions on Signal Processing*, vol. 66, no. 10, pp. 2746–2758, 2018.
- [14] C. You, B. Zheng, and R. Zhang, "Channel estimation and passive beamforming for intelligent reflecting surface: Discrete phase shift and progressive refinement," *IEEE Journal on Selected Areas in Communications*, vol. 38, no. 11, pp. 2604–2620, 2020.
- [15] Q. Wu and R. Zhang, "Towards smart and reconfigurable environment: Intelligent reflecting surface aided wireless network," *IEEE Communications Magazine*, vol. 58, no. 1, pp. 106–112, 2019.
- [16] S. Bi and Y. J. Zhang, "Computation rate maximization for wireless powered mobile-edge computing with binary computation offloading," *IEEE Transactions on Wireless Communications*, vol. 17, no. 6, pp. 4177–4190, 2018.
- [17] S. Zhang, H. Gu, K. Chi, L. Huang, K. Yu, and S. Mumtaz, "DRL-based partial offloading for maximizing sum computation rate of wireless powered mobile edge computing network," *IEEE Transactions on Wireless Communications*, vol. 21, no. 12, pp. 10934–10948, 2022.
- [18] L. Huang, S. Bi, and Y.-J. A. Zhang, "Deep reinforcement learning for online computation offloading in wireless powered mobile-edge computing networks," *IEEE Transactions on Mobile Computing*, vol. 19, no. 11, pp. 2581–2593, 2020.
- [19] V. D. P. Souto, R. D. Souza, B. F. Uchoa-Filho, A. Li, and Y. Li, "Beamforming optimization for intelligent reflecting surfaces without CSI," *IEEE Wireless Communications Letters*, vol. 9, no. 9, pp. 1476–1480, sep 2020.
- [20] Y. Yang, S. Zhang, and R. Zhang, "Irs-enhanced ofdm: Power allocation and passive array optimization," in *2019 IEEE Global Communications Conference (GLOBECOM)*, 2019, pp. 1–6.
- [21] S. Zhang and R. Zhang, "Capacity characterization for intelligent reflecting surface aided mimo communication," *IEEE Journal on Selected Areas in Communications*, vol. 38, no. 8, pp. 1823–1838, 2020.
- [22] Y. Yang, Y. Gong, and Y.-C. Wu, "Intelligent-reflecting-surface-aided mobile edge computing with binary offloading: Energy minimization for iot devices," *IEEE Internet of Things Journal*, vol. 9, no. 15, pp. 12973–12983, 2022.
- [23] C. Sun, W. Ni, Z. Bu, and X. Wang, "Energy minimization for intelligent reflecting surface-assisted mobile edge computing," *IEEE Transactions on Wireless Communications*, vol. 21, no. 8, pp. 6329–6344, 2022.
- [24] G. Chen, Q. Wu, R. Liu, J. Wu, and C. Fang, "Irs aided mec systems with binary offloading: A unified framework for dynamic irs beamforming," *IEEE Journal on Selected Areas in Communications*, vol. 41, no. 2, pp. 349–365, 2023.
- [25] M. Zeng, R. Du, V. Fodor, and C. Fischione, "Computation rate maximization for wireless powered mobile edge computing with noma," in *2019 IEEE 20th International Symposium on "A World of Wireless, Mobile and Multimedia Networks" (WoWMoM)*, 2019, pp. 1–9.
- [26] P.-Q. Huang, Y. Wang, K. Wang, and Q. Zhang, "Combining lyapunov optimization with evolutionary transfer optimization for long-term energy minimization in irs-aided communications," *IEEE Transactions on Cybernetics*, vol. 53, no. 4, pp. 2647–2657, 2023.
- [27] Q. Wu and R. Zhang, "Intelligent reflecting surface enhanced wireless network via joint active and passive beamforming," *IEEE Transactions on Wireless Communications*, vol. 18, no. 11, pp. 5394–5409, 2019.
- [28] S. Mao, S. Leng, K. Yang, X. Huang, and Q. Zhao, "Fair energy-efficient scheduling in wireless powered full-duplex mobile-edge computing systems," in *GLOBECOM 2017 - 2017 IEEE Global Communications Conference*. IEEE, dec 2017, pp. 1–6.
- [29] L. Shi, Y. Ye, X. Chu, and G. Lu, "Computation energy efficiency maximization for a noma-based wpt-mec network," *IEEE Internet of Things Journal*, vol. 8, no. 13, pp. 10731–10744, 2021.
- [30] L. Ji and S. Guo, "Energy-efficient cooperative resource allocation

in wireless powered mobile edge computing," *IEEE Internet of Things Journal*, vol. 6, no. 3, pp. 4744–4754, jun 2019.

- [31] F. Zhou and R. Q. Hu, "Computation efficiency maximization in wireless-powered mobile edge computing networks," *IEEE Transactions on Wireless Communications*, vol. 19, no. 5, pp. 3170–3184, 2020.
- [32] S. Mao, S. Leng, K. Yang, X. Huang, and Q. Zhao, "Fair energy-efficient scheduling in wireless powered full-duplex mobile-edge computing systems," in *GLOBECOM 2017 - 2017 IEEE Global Communications Conference*, 2017, pp. 1–6.
- [33] H. Lim and T. Hwang, "Energy-efficient computing for wireless powered mobile edge computing systems," in *2019 IEEE 90th Vehicular Technology Conference (VTC2019-Fall)*, 2019, pp. 1–5.
- [34] W. Zheng and L. Yan, "Latency minimization for IRS-assisted mobile edge computing networks," *Physical Communication*, vol. 53, p. 101768, aug 2022.
- [35] H. Guo, Y. C. Liang, J. Chen, and E. G. Larsson, "Weighted sum-rate maximization for reconfigurable intelligent surface aided wireless networks," *IEEE Transactions on Wireless Communications*, vol. 19, no. 5, pp. 3064–3076, 2020.
- [36] Y. Wang, S. Min, X. Wang, W. Liang, and J. Li, "Mobile-edge computing: Partial computation offloading using dynamic voltage scaling," *IEEE Transactions on Communications*, vol. 64, no. 10, pp. 4268–4282, 2016.
- [37] T. D. Burd and R. W. Brodersen, "Processor design for portable systems," *J VLSI Sign Process Syst Sign Image Video Technol*, vol. 13, no. 2-3, pp. 203–221, aug 1996.
- [38] M. Xu and R. Buyya, "Managing renewable energy and carbon footprint in multi-cloud computing environments," *Journal of Parallel and Distributed Computing*, vol. 135, pp. 191–202, 2020.
- [39] W. Dinkelbach, "On nonlinear fractional programming," *Management Science*, vol. 13, no. 7, pp. 492–498, mar 1967.
- [40] K. Shen and W. Yu, "Fractional programming for communication systems—part II: Uplink scheduling via matching," *IEEE Transactions on Signal Processing*, vol. 66, no. 10, pp. 2631–2644, 2018.
- [41] L. Shi, Y. Ye, X. Chu, and G. Lu, "Computation energy efficiency maximization for a NOMA-based WPT-MEC network," *IEEE Internet of Things Journal*, vol. 8, no. 13, pp. 10731–10744, jul 2021.



Sukhpal Singh Gill is a Lecturer (Assistant Professor) in Cloud Computing at the School of Electronic Engineering and Computer Science, Queen Mary University of London, UK. Prior to his present stint, Dr. Gill has held positions as a Research Associate at the Lancaster University, UK and also as a Postdoctoral Research Fellow at CLOUDS Laboratory, The University of Melbourne, Australia. Dr. Gill is serving as an Associate Editor in IEEE IoT, Wiley SPE, Elsevier IoT, Wiley ETT and IET Networks Journal. He has published in prominent international journals and conferences such as IEEE TCC, IEEE TSC, IEEE TII, IEEE TNSM, IEEE IoT Journal, Elsevier JSS/FGCS, IEEE/ACM UCC and IEEE CCGRID. His research interests include Cloud Computing, Fog Computing, IoT and Energy Efficiency. For further information, please visit: <http://www.ssgill.me>.



Huaming Wu (Senior Member, IEEE) received the B.E. and M.S. degrees from Harbin Institute of Technology, China in 2009 and 2011, respectively, both in electrical engineering. He received the Ph.D. degree with the highest honor in computer science at Freie Universität Berlin, Germany in 2015. He is currently an associate professor at the Center for Applied Mathematics, Tianjin University, China. His research interests include mobile cloud computing, edge computing, internet of things, deep learning, complex

network, and DNA storage.



Junhui Du received the BSc degree in mathematics from the Nanjing University of Information Science Technology, China, in 2021. He is currently working toward the master's degree in mathematics with the Center for Applied Mathematics, Tianjin University, Tianjin, China. His research interests include Internet of Things, deep learning, and mobile edge computing.



Minxian Xu (Member, IEEE) is currently an Associate Professor at Shenzhen Institutes of Advanced Technology, Chinese Academy of Sciences. He received the BSc degree in 2012 and the MSc degree in 2015, both in software engineering from University of Electronic Science and Technology of China. He obtained his Ph.D. degree from the University of Melbourne in 2019. His research interests include resource scheduling and optimization in cloud computing. He has co-authored 40+ peer-reviewed papers

published in prominent international journals and conferences, such as ACM CSUR, ACM TOIT, IEEE TSUSC, IEEE TCC, IEEE TASE, IEEE TGCN, JPDC, JSS and IC SOC. His Ph.D. Thesis was awarded the 2019 IEEE TCSC Outstanding Ph.D. Dissertation Award. More information can be found at: minxianxu.info.