

Human-in-the-Loop Wireless Communications: Machine Learning and Brain-Aware Resource Management

Ali Taleb Zadeh Kasgari¹, Walid Saad¹, and Mérouane Debbah²

¹ Wireless@VT, Electrical and Computer Engineering Department, Virginia Tech, VA, USA,

Emails: {alitek, walids}@vt.edu.

² Mathematical and Algorithmic Sciences Lab, Huawei France R&D, Paris, France, and CentraleSupélec,

Université Paris-Saclay, Gif-sur-Yvette, France, Email: merouane.debbah@huawei.com.

Abstract

Human-centric applications such as virtual reality and immersive gaming are central to future wireless networks. Common features of such services include: a) their dependence on the human user's behavior and state, and b) their need for more network resources compared to conventional applications. To successfully deploy such applications over wireless networks, the network must be made cognizant of not only the quality-of-service (QoS) needs of the applications, but also of the perceptions of the *human users* on this QoS. In this paper, by explicitly modeling the limitations of the human brain, a concrete measure for the delay perception of human users is introduced. Then, a learning method, called probability distribution identification, is developed to find a probabilistic model for this delay perception based on the brain features of a human user. Given the learned model for the delay perception of the human brain, a brain-aware resource management algorithm based on Lyapunov optimization is proposed for allocating radio resources to human users while minimizing the transmit power and taking into account the reliability of both machine type devices and human users. Then, a closed-form relationship between the reliability measure and wireless physical layer metrics of the network is derived. Simulation results show that a brain-aware approach can yield savings of up to 78% in power compared to the system that only considers QoS metrics. The results also show that, compared with QoS-aware, brain-unaware systems, the brain-aware approach can save substantially more power in low-latency systems.

A preliminary version of this work appeared in the proceedings of the 51th Asilomar Conference on Signals, Systems and Computers, Pacific Grove, CA, USA [1].

This research was supported by the U.S. National Science Foundation under Grants CNS-1460316 and IIS-1633363.

I. INTRODUCTION

The next generation of wireless services is expected to be highly human centric. Examples include virtual reality and interactive/immersive gaming [2]–[4]. To cope with the quality-of-service (QoS) needs of such human-centric applications, in terms of data rate and ultra-low latency, wireless networks must exploit substantially more radio resources by leveraging heterogeneous spectrum bands [5]. Although allocating heterogeneous spectrum resources can potentially increase the raw QoS, given the human-centric nature of emerging applications, their users may not be able to perceive the improved QoS, due to human factors such as the cognitive limitations of the brain [6]. Indeed, many empirical studies (anecdotal and otherwise) have shown that the limitations on the human brain can be translated into a limitation on how wireless users translate QoS into actual quality-of-experience (QoE) [7]–[9]. For example, the human brain may not be able to perceive any difference between videos transmitted with different QoS (e.g., rates or delays) [9], [10].

Hence, in order to deploy these services over wireless networks, such as 5G cellular systems, there is a need to enable the system to be strongly cognizant of the human user in the loop. In particular, to deliver immersive, human-centric services, the network must tailor the usage and optimization of wireless resources to the intrinsic features of its human users such as their behavior and brain processing limitations. By doing so, the network can potentially save resources, accommodate more users, and provide a more realistic QoE to its users. Moreover, the saved resources can be used to accommodate emerging applications in wireless networks such as drone communications [11], [12] and autonomous driving [13]–[16].

Developing resource management mechanisms that can cater to intrinsic needs of wireless users and their context (e.g., device features or social metrics) has recently been studied in [5], [17]–[24]. In [17], a context-aware scheduling algorithm for 5G systems is proposed. This algorithm exploits the context information of user equipments (UEs), such as battery level, to save energy in the system while satisfying the QoS requirements of users. The authors in [18], proposed a user-centric resource allocation framework for ultra-dense heterogeneous networks. Context-aware resource allocation for heterogeneous cellular networks is also studied in [5], [19], and [20]. In [5], a novel approach to context-aware resource allocation in small cell networks is introduced. Both wireless physical layer metrics and the social ties of human users are exploited in [5] to allocate wireless resource blocks. Proactive caching using context information from

social networks is studied in [21]. The results in [21] show that such a socially-aware caching technique reduces the peak traffic in 5G networks. Other context-aware resource allocation algorithms are also studied in [22], [25], and [24]. However, despite this surge in literature on context-aware networking [5], [17]–[22], [24], and [25] this prior art is still reliant on device-level features. Moreover, the works in [5], [17]–[22], [24], and [25] are agnostic to the human users and their features (e.g., brain limitation or behavior). Hence, adapting these existing approaches can waste network resources due to the potential allocation of more resources to human users that cannot perceive the associated QoS gains, due to cognitive brain limitations.

A general framework for modeling the intelligence of communication systems which serve humans is proposed in [26]. The author defines intelligence in terms of predicting and serving human demands in advance. However, the work in [26] does not account for the cognitive limitations of a human brain. Moreover, demand prediction, as done in [26], will not be sufficient to capture the full spectrum of the human user limitations and behavior. By being aware of brain limitations of each user, the network can provide a unique experience for each user and optimize its performance. For example, an increase in the delay of a wireless system may have different effects on the QoE perceived by different human users. In particular, such different delay perceptions can potentially be exploited by the cellular network to minimize power consumption and reduce the amount of wasted resources. To our best knowledge, no existing work has studied the impact of such disparate brain delay perceptions on wireless resource allocation. systems. Furthermore, none of the prior studies on systems with humans-in-the-loop in other fields [27] and [28] have analyzed the human brain limitations.

The main contribution of the paper is a *novel brain-aware learning and resource management framework* that explicitly factors in the brain state of human users during resource allocation in a cellular network. In particular, we formulate the brain-aware resource allocation problem using a joint learning and optimization framework. First, we propose a learning algorithm to identify the delay perceptions of a human brain. This learning algorithm employs both supervised and unsupervised learning to identify the brain limitations and also creates a statistical model for these limitations based on Gaussian mixture models. Then, using Lyapunov optimization, we address the resource allocation problem with time varying QoS requirements that captures the learned delay perception. Using this approach, the network can allocate radio resources to human users while considering the reliability of both machine type devices and human users. We then identify a closed-form relationship between system reliability and wireless physical layer metrics

and derive a closed-form expression for the reliability as a function of the human brain's delay perception. Simulation results using real data show that the proposed brain-aware approach can substantially save power in the network while preserving the reliability of the users, particularly in low latency applications. In particular, the results show that the proposed brain-aware approach can yield power savings of up to 78% compared to a conventional, brain-unaware system.

The rest of the paper is organized as follows. Section II introduces the system model. Sections III and IV present the proposed learning algorithm and resource allocation framework, respectively. Section V presents the simulation results and conclusions are drawn in Section VI.

II. SYSTEM MODEL AND PROBLEM FORMULATION

Consider the downlink of a cellular network with humans-in-the-loop having a single base station (BS) serving a set $\mathcal{H}$ of N human users with their UEs and a set $\mathcal{M}$ of M machine type devices (MTDs). We assume that each human user uses one UE. Each UE or MTD can have a different application with different QoS requirements such as sending a command to an actuator (for an MTD) or playing a 3D interactive game (for a UE). We consider a time-slotted system, and define $\mathcal{K}$ as the set of K resource blocks (RBs). In our model, the packets associated with user $i \in \mathcal{H} \cup \mathcal{M}$ arrive at the BS according to independent Poisson processes with rate $a_i(t)$. The lengths $l_i, \forall i \in \mathcal{H} \cup \mathcal{M}$ of the packets follow an exponential distribution. Hence, each user's buffer at the BS will follow an M/M/1 queuing model. The total queuing and transmission delay of each user i is $D_i(t) = q_i(t) + \frac{l_i}{r_i(t)}$, where $q_i(t)$ is the queuing delay. The data rate for each user is given by:

$$r_i(t) = B \sum_{j=1}^K \rho_{ij}(t) \log_2 \left(1 + \frac{p_{ij}(t)h_{ij}(t)}{\sigma^2} \right), \quad (1)$$

where $p_{ij}(t)$ is the transmit power between the BS and user i over RB j at time t and $h_{ij}(t)$ is the time-varying Rayleigh fading channel gain. In (1), $\rho_{ij}(t) = 1$ if RB j is allocated to user i at time slot t , and $\rho_{ij}(t) = 0$, otherwise. B is the bandwidth of each RB. σ^2 is the noise power which is defined as the power spectral density of the noise multiplied by the bandwidth B .

We define $\beta_i(t)$ as the delay perception threshold for any user $i \in \mathcal{H} \cup \mathcal{M}$ at time t . If the delay decreases below the threshold $\beta_i(t)$, the user will not be able to discern the change in service quality. We use the concept of delay perception $\beta_i(t)$ to measure time varying delay requirements of UEs and MTDs. Since the delay perception for MTDs is constant, hereinafter, for simplicity, we use $\beta_i(t)$ to exclusively denote the delay perception of human users, i.e., $\beta_i(t), \forall i \in \mathcal{H}$, unless

mentioned otherwise. This delay perception can be affected by multiple sources pertaining to the human brain such as context, human attention, human fatigue or cognitive abilities and is determined by measuring the capabilities of the human brain at each time slot using machine learning methods. *By explicitly accounting for the cognitive limitations of the human brain, the BS can better allocate resources to the users that need it, when they can actually use it.* This is in contrast to conventional brain-agnostic networks [5], [26] in which resources may be wasted, as they are allocated only based on application QoS without being aware on whether the human user can indeed process the actual application's QoS target.

We pose this resource allocation problem as a power minimization problem that is subject to a brain-aware QoS constraint on the latency:

$$\min_{\boldsymbol{\rho}(t), \mathbf{P}(t)} \sum_{j \in \mathcal{K}} \left[\sum_{i \in \mathcal{H}} \bar{P}_i^j + \sum_{i \in \mathcal{M}} \bar{P}_i^j \right], \quad (2a)$$

$$\text{s.t.} \quad \Pr\{D_i(t) \geq D_i^{\max}(\beta_i(t))\} \leq \epsilon_i(\beta_i(t)), \quad \forall i \in \mathcal{H} \cup \mathcal{M}, \quad (2b)$$

$$p_{ij}(t) \geq 0, \quad \rho_{ij}(t) \in \{0, 1\} \quad \forall i \in \mathcal{H} \cup \mathcal{M}, j \in \mathcal{K}, \quad (2c)$$

$$\sum_{i \in \mathcal{H} \cup \mathcal{M}} \rho_{ij}(t) = 1, \quad \forall j \in \mathcal{K}, \quad (2d)$$

where $\boldsymbol{\rho}(t)$ is an $(M + N) \times K$ matrix having each element $\rho_{ij}(t)$. $\mathbf{P}(t)$ is an $(M + N) \times K$ matrix with each element $p_{i,j}(t)$ representing the instantaneous power allocated to user i on RB j . The term $\bar{P}_i^j = \lim_{t \rightarrow \infty} \frac{1}{t} \sum_{\tau=0}^{t-1} \rho_{ij}(\tau) p_{ij}(\tau)$ is the time average of the power allocated to user i on RB j . $D_i(t)$ incorporates both transmission and queuing delays. $D_i^{\max}(\beta_i(t))$ is the maximum tolerable delay. This delay depends on $\beta_i(t)$ because changes in human delay perception will change maximum tolerable delay for human users. $\epsilon_i(\beta_i(t))$ in equation (2b) denotes the maximum probability of the packet delay exceeding $D_i^{\max}(\beta_i(t))$. Hence, we can define $1 - \epsilon_i(\beta_i(t))$ as the reliability of the user i . We define *reliability* as the proportion of time during which the delay of a given user does not exceed a threshold. For notational convenience, hereinafter, we use the terms $D_i^{\max}(\beta_i(t))$ and $D_i^{\max}$ interchangeably. Constraint (2b) takes into account the packet size and the rate of the application implicitly in addition to the maximum tolerable delay and reliability. From a resource allocation perspective, we can consider any application using (2b).

The key difference between our problem formulation and conventional RB allocation problems [29] is seen in the QoS delay requirement in (2b). Constraint (2b) is with respect to two random processes $D_i(t)$ and $\beta_i(t)$. In (2b), the network explicitly accounts for the human brain's (and

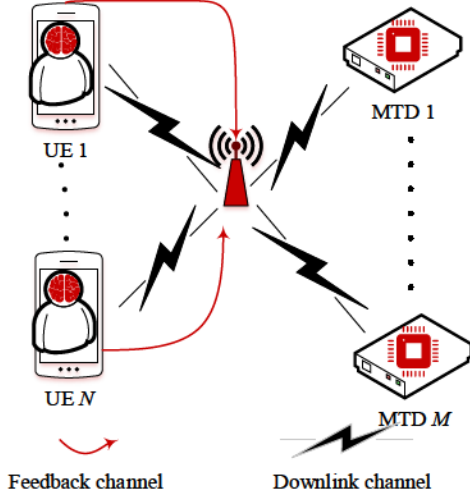


Figure 1: Illustration of the system model.

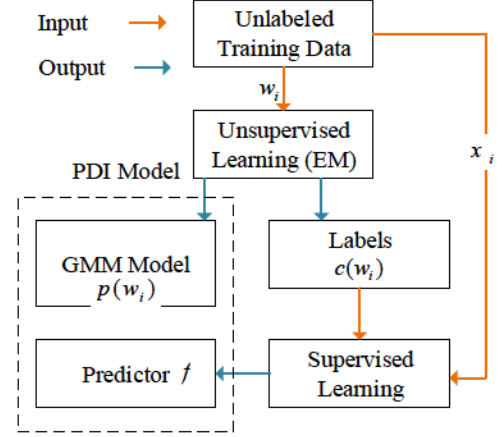


Figure 2: Graphical representation of building a PDI model.

the MTDs') delay needs. By taking into account the features of the brain of the human UEs, the network can avoid wasting resources. This waste of resources can stem from allocating more power to a UE, solely based on the application QoS, while ignoring how the brain of the human carrying the UE perceives this QoS. Clearly, ignoring this human perception can lead to inefficient resource management.

We propose a machine learning algorithm to identify the human brain delay perception $\beta_i(t)$. Each human user has d features, (e.g., age, occupation, location) assumed to be known to the BS. This time-varying feature vector is denoted by $x_i(t) \in \mathbb{R}^d$. We develop a learning algorithm to build a model that maps these features to $\beta_i(t)$ for each user. We then show that being aware of $\beta_i(t)$ can help the resource allocation algorithm to save a significant amount of resources for low-latency systems. We assume that the BS has access to the user features $x_i(t)$. In practice, the BS can collect such data whenever a given user registers in the network or by using the sensors of a user's mobile device. The system model is shown in Fig. 1. Also, Table I provides a list of our main parameters and notations.

To find the mapping $\beta_i(t) = f(x_i(t))$ between human features $x_i(t)$ and the delay perception of the brain, we introduce a novel supervised learning mechanism called the *probability distribution identification (PDI) method*. Here, function $f(\cdot)$ shows this mapping. Since reliability is a key factor in a communication system, we need a supervised learning algorithm that not only predicts

Table I: List of notations.

Notation	Description	Notation	Description
p_{ij}	Power of user i on RB j	ρ_{ij}	RB j allocation indicator to user i
D_i	Packet delay for user i	L	Total number of clusters, brain modes
$D_i^{\max}(\cdot)$	Target delay for user	$\epsilon_i(\cdot)$	target reliability of user i
$\beta_i(t)$	Delay perception of user i at time t	V	Balancing parameter for reliability constraint
$h_{ij}(t)$	Time-varying Rayleigh fading channel gain	λ	Lagrange multiplier
B	RB bandwidth	σ^2	Noise power
N	Number of UEs	$k(j)$	Optimal allocation for RB j
K	Number of available RBs	Π	Projection operator
$a_i(t)$	Arrival rate of user i in time t	ζ_i	Subgradient of element i of λ_b
$F_i(t)$	Virtual queue for user i at time t	Σ_k	Covariance matrix of mode k of human brain
n	# of training users,	$\psi(\cdot)$	PDF of multivariate normal
$D_i^{\max}$	Target end-to-end latency for user i	μ_k	Mean vector of mode k of human brain
$\mathcal{M}$	set of MTDs	$\mathcal{H}$	Set of UEs
$\mathcal{K}$	Set of RBs	$\chi_{d+1}^2(\cdot)$	chi-square distribution function
π_k	Mixture weight of mode k	$\mathbf{z}$	Mode indicator vector
$\mathcal{L}$	Likelihood function	d	# of features for each user
f	Supervised learning model	$c(\mathbf{w}_i), c_i$	Cluster (mode) number of $\mathbf{w}_i$
$\mathbf{y}$	Set of labels for supervised learning	$\xi(\cdot)$	0–1 loss function
$D_i^{\min}(\epsilon')$	Effective delay of human user i	$Q_d(\gamma)$	Quantile function of chi-square distribution with d DOF
$\mathcal{S}$	Reliability event of the system	$\chi_d(\cdot)$	Chi-square function with d DOF
$\mathcal{L}_a$	Lagrange dual function	$\mathbf{e}_j$	Unit vector at j

$\beta_i(t)$ as a function of $\mathbf{x}_i(t)$, but also gives a measure of reliability for this prediction. This *measure of reliability* is one of the key advantages of PDI learning over other supervised learning methods [30], [31]. Although conventional methods such as neural networks can be used to approximate the continuous function $f(\cdot)$ [31], these methods cannot quantify the reliability of this prediction. The reliability of predictions is defined as the probability that the prediction of $\beta_i(t)$ lies within a certain range of the true values for $\beta_i(t)$.

The PDI method can find the distribution of the prediction values. Although many existing supervised learning methods can build a model for predicting an output based on a given input, they fail to find a statistical model for this prediction. In addition, by using the statistical model for the predictions resulting from our proposed PDI, the system designer can better design the system based on desired reliability.

As discussed in [32], the delay perception of a human brain typically follows a multi-modal distribution. As a result, we design the proposed PDI approach to capture such a model and find the different modes of a human brain. Then, using the distribution of the brain delay, the PDI approach can find the *effective delay* of the human brain. This effective delay determines relationship between $\beta_i(t)$ and $\mathbf{x}_i(t)$ along with its reliability.

A. Building a PDI model

Consider a dataset $\{\mathbf{x}_1(t), \dots, \mathbf{x}_n(t)\}$, where $\mathbf{x}_i(t) \in \mathbb{R}^d$ is one sample data vector. The elements of $\mathbf{x}_i(t)$ are user features which can be both categorical (such as gender) and numerical (such as age). For each input vector $\mathbf{x}_i(t)$, we have a corresponding output value of delay perception $\beta_i(t)$. This data can be collected using experiments or surveys such as those in [33]. Since we can remove time dependency of the data using time-series techniques such as in [34], hereinafter, we use $\mathbf{x}$ instead of $\mathbf{x}(t)$. Although we omit the time dependency from $\mathbf{x}_i(t)$ for the training process, it is still implicitly a function of time. This dataset can be represented by a matrix $\mathbf{X} \in \mathbb{R}^{n \times d}$, where $\mathbf{x}_i^T$ is row i of $\mathbf{X}$. Using PDI, we first create an $n \times (d+1)$ dataset matrix:

$$\mathbf{W} = [\mathbf{X} \parallel \boldsymbol{\beta}] = \begin{bmatrix} \mathbf{w}_1^T \\ \vdots \\ \mathbf{w}_n^T \end{bmatrix} = \begin{bmatrix} \mathbf{x}_1^T & \beta_1(t) \\ \vdots & \vdots \\ \mathbf{x}_n^T & \beta_n(t) \end{bmatrix}, \quad (3)$$

where $\mathbf{w}_i \in \mathbb{R}^{d+1}$ is a vector of the delay perception $\beta_i(t)$ and d other correlated features of the human brain. First, in the unsupervised learning step, we fit a Gaussian mixture model (GMM) to our dataset using the expectation-maximization (EM) algorithm [35] to obtain $p(\mathbf{x}_i, \beta_i(t))$. After finding $p(\mathbf{x}_i, \beta_i(t))$, we are able to cluster the data samples and find m brain modes in the data. Then, each data vector $\mathbf{x}_i$ is labeled based on its cluster number c_i so that each $\mathbf{x}_i$, $i = 1, \dots, n$ has a label in the cluster set $c_i \in \mathcal{C} = \{1, \dots, L\}$. Using this method we have a labeled dataset which can be used for the supervised learning. *These cluster numbers will correspond to the modes of the human brain that determine its effective delay perception.*

Next, we describe the Gaussian mixture model that we use as underlying model for the human brain. We use GMM for clustering (unsupervised learning) and statistical modeling of human brain because of its multimodal structure which resembles human brain activities [36], its scalability, and its robustness and stability under high-noise levels compared to nonparametric methods.

A multi-modal stochastic model is assumed for the brain features $\mathbf{w}_i$ for user i . The proposed distribution for $\mathbf{w}_i$ is given by [30]:

$$p(\mathbf{w}_i) = \sum_z p(\mathbf{z})p(\mathbf{w}_i|\mathbf{z}) = \sum_{k=1}^L \pi_k \psi(\mathbf{w}_i|\boldsymbol{\mu}_k, \boldsymbol{\Sigma}_k), \quad (4)$$

where $\psi(\mathbf{w}_i|\boldsymbol{\mu}_k, \boldsymbol{\Sigma}_k)$ is the probability density function for a multivariate normal distribution

with mean vector $\boldsymbol{\mu}_k$ and covariance matrix $\boldsymbol{\Sigma}_k$. $\boldsymbol{\Sigma}_k$ and $\boldsymbol{\mu}_k$ represent the covariance matrix and mean vector for mode k of the human brain, respectively. $\mathbf{z}$ is a binary random vector, in which a particular element z_k is equal to 1 and all other elements are 0. $\mathbf{z}$ essentially indicates which mode is activated in the GMM, and π_k is defined as $p(z_k = 1) = \pi_k$. L is the total number of modes in the GMM. The human brain will be in mode k with probability π_k , and its features are generated using a multivariate normal distribution with mean and covariance $\boldsymbol{\mu}_k$ and $\boldsymbol{\Sigma}_k$, respectively. The posterior probability, i.e. *responsibility*, for mode k will be:

$$r_i(z_k) = \frac{\pi_k \psi(\mathbf{w}_i | \boldsymbol{\mu}_k, \boldsymbol{\Sigma}_k)}{\sum_{j=1}^L \pi_j \psi(\mathbf{w}_i | \boldsymbol{\mu}_j, \boldsymbol{\Sigma}_j)}. \quad (5)$$

This responsibility can be used for clustering the data as well. After fitting the GMM on the dataset, we can find the mode with highest responsibility for each data point and assign the data to this mode. The EM algorithm is used to find $\boldsymbol{\mu}_k$, $\boldsymbol{\Sigma}_k$, and π_k for all $k = 1, \dots, L$, based on the real-time human brain behavior [35]. The log likelihood function for our dataset can be written as:

$$\ln \mathcal{L}(\boldsymbol{\Sigma}, \boldsymbol{\mu}, \boldsymbol{\pi} | \mathbf{w}) = \ln \sum_i p(\mathbf{w}_i | \boldsymbol{\Sigma}, \boldsymbol{\mu}, \boldsymbol{\pi}) = \sum_i \ln \sum_{k=1}^L \pi_k \psi(\mathbf{w}_i | \boldsymbol{\mu}_k, \boldsymbol{\Sigma}_k). \quad (6)$$

The likelihood function in (6) has singularities and, hence, it is infeasible to find parameters π_k , $\boldsymbol{\Sigma}_k$, and $\boldsymbol{\mu}_k$. The EM algorithm is proposed in [37] to maximize the likelihood function for a Gaussian mixture model. In the EM algorithm, we first initialize $\boldsymbol{\Sigma}_k$, $\boldsymbol{\mu}_k$, and π_k randomly. Next, we find the responsibility for each mode using (5). Then, we reestimate parameters using current responsibilities. Finally, the likelihood in (6) is maximized with respect to $\boldsymbol{\Sigma}_k$, $\boldsymbol{\mu}_k$, and π_k .

As a result of the EM algorithm (unsupervised learning), we now have a GMM of our dataset matrix $\mathbf{W}$. Based on this GMM, the data will be labeled (clustered) as follows. For each data point $\mathbf{w}_i$, the most probable mode is assigned as the label of this data, i.e.,

$$c_i = c(\mathbf{w}_i) = \arg \max_k p(z_k = 1 | \mathbf{w}_i) = \arg \max_k r_i(z_k). \quad (7)$$

In (7), we assign the most likely cluster to each data point $\mathbf{w}_i$. After building a GMM model using unsupervised learning on the $\mathbf{W}$ dataset to obtain target vector $\mathbf{y} = [c(\mathbf{w}_1) \ \dots \ c(\mathbf{w}_n)]^T$, we use the pair $\{\mathbf{X}, \mathbf{y}\}$ to train a supervised learning model. Thus, the output of the unsupervised learning step $\mathbf{y}$ is used for training the supervised learning model. Then, during the supervised learning step, we train a classifier so that it can find the mode c_i using the human features $\mathbf{x}_i$ as input. Given the data matrix $\mathbf{X}$ and the output vector $\mathbf{y}$, this supervised learning builds us a

Algorithm 1 Building PDI model

Input: $w_i = [x_i \ \beta_i]$, $i = 1, \dots, n$

Output: $f, \pi_k, \mu_k, \Sigma_k, \quad k = 1 = \dots, L$

Unsupervised Learning :

- 1: Apply EM algorithm to w_i , find $\pi_k, \mu_k, \Sigma_k, r_i(z_k) \quad k = 1, \dots, L, \quad i = 1 \dots, n$.
- 2: **for** $i = 1 = \dots, n$ **do**
- 3: find $c(w_i)$ using (7).
- 4: **end for**
- 5: Pass $\mathbf{y} = [c(w_1) \ \dots \ c(w_n)]^T$ to the supervised learning algorithm

Supervised Learning :

- 6: Find f using (8).
 - 7: **return** $f, \pi_k, \mu_k, \Sigma_k, \quad k = 1 = \dots, L$
-

model f such that $c_i = f(x_i)$, where

$$f = \arg \min_{\hat{f}} \sum_{i=1}^n \xi \left(c(w_i), \hat{f}(x_i) \right), \quad (8)$$

where $\xi(\cdot)$ is a 0-1 loss function [38, Equation 7.5]. f is a function that is approximated using a set of points (x_i, c_i) and determines the relationship between the features of a user and its cluster. To overcome overfitting, we use the elbow method for finding the optimal number of clusters in PDI, as discussed in Section IV. After approximating f , given each human user's feature vector x_i , we find the modes c_i using model f . Algorithm 1 summarizes building a PDI learning model. Also, a graphical representation for building a PDI model is shown in Fig. 2.

B. Deployment of the PDI learning model

In this subsection, we bound $D_i^{\max}(\beta_i(t))$ based on its features x_i . Note that the deployment and training of the PDI learning method are separated from each other. The deployment part only uses the model generated by the training part. Now that the system can identify the human users' modes, we need to find a relationship between a human user's mode and the probabilistic model of its delay perception by defining the concept of effective delay.

Definition 1. Given the statistical model for human delay perception $\beta_i(t)$, $D_i^{\min}(\epsilon')$ is the *effective delay* for human user i that satisfies:

$$\Pr\{\beta_i(t) < D_i^{\min}(\epsilon')\} < \epsilon'. \quad (9)$$

To find the effective delay for human user i , we first find the probability that the delay perception of human user i is less than a threshold $D_i^{\min}(\epsilon')$. In other words, we want to find the relation between ϵ' and $D_i^{\min}(\epsilon')$ in (9). The concept of effective delay is defined using the fact that delays less than $D_i^{\min}(\epsilon')$ cannot be sensed by a human with $(1 - \epsilon')$ certainty. The relation

Algorithm 2 Deploying PDI model for finding brain delay perception

Input: $x_i, i = 1, \dots, N$
Output: $\mu_i, \Sigma_i, i = 1, \dots, N$

- 1: **for** $i = 1 = \dots, N$ **do**
 - 2: find the cluster for x_i : $k_i = f(x_i)$ and hence μ_{k_i}, Σ_{k_i} .
 - 3: find function $D_i^{\min}(\epsilon')$ using (13)
 - 4: using target reliability find ϵ' and then using function $D_i^{\min}(\epsilon')$ find $D_i^{\min}$.
 - 5: set the $D_i^{\max}(\beta_i(t)) = D_i^{\min}$, and using (16) and ϵ' find $\epsilon(\beta_i(t))$
 - 6: **end for**
 - 7: **return** $D_i^{\max}(\beta_i(t))$ and $\epsilon(\beta_i(t))$ for $i = 1 = \dots, N$
-

between ϵ' and $D_i^{\min}(\epsilon')$ in (9) is found in Theorem 1. For notational simplicity, hereinafter, we use $D_i^{\min}$ instead of $D_i^{\min}(\epsilon')$.

Theorem 1. If brain mode k is identified for user i , then its delay perception will be bounded as follows:

$$\Pr\left\{|\beta_i(t) - \mu_k(d+1)| < \sqrt{Q_{d+1}(\gamma) \mathbf{e}_{d+1}^T \Sigma_k \mathbf{e}_{d+1}}\right\} > \gamma, \quad (10)$$

where Σ_k and $\mu_k(d+1)$ represent, respectively, the covariance matrix and the $(d+1)$ th element of the mean vector of the identified brain mode k . $Q_d(\gamma)$ is the quantile function of chi-square distribution with d degrees of freedom, and is defined as

$$Q_{d+1}(\gamma) = \inf \left\{ x \in \mathbb{R} \mid \gamma \leq \int_0^x \chi_{d+1}^2(u) du \right\}, \quad (11)$$

and $\mathbf{e}_j$ is a unit vector in $\mathbb{R}^{d+1}$, whose j th element is 1 and all other elements are zero. d is number of features used for learning. $\chi_{d+1}^2(x)$ is the probability density function of a chi-square random variable with $d+1$ degrees of freedom.

Proof: See Appendix A. ■

Note that the bound in (10) is different from finding a bound using marginal distributions, as it is based on high probability density areas. The use of a marginal distribution is not possible here, since the Gaussian assumption is only valid locally around the mean. Also, in case of data classification error which mostly happens when the data is located in the overlapping area between clusters, the bound in (10) will be either more conservative than the actual bound for delay perception or it will not change significantly compared to the actual bound. As seen from Theorem 1, in addition to the delay perception element $\mu_k(d+1)$, the only other parameter that affects the delay is $\mathbf{e}_{d+1}^T \Sigma_k \mathbf{e}_{d+1}$, which is the $(d+1)$ th diagonal element of Σ_k , which is not assumed to be diagonal. Fig.3 shows the relationship between $D_i^{\min}$ and GMM random data

generated from a Gaussian mixture model. Fig.3 shows that, after finding the GMM for the dataset, one can find the predictive coverage of each Gaussian distribution and, then, we can determine the probability with which $\beta_i(t)$ for a user i will be higher than a threshold $D_i^{\min}$. In order to find the effective delay for human user i , we first find the probability with which the delay perception for human user i will be less than a threshold $D_i^{\min}$. In other words, we will find the relationship between ϵ and $D_i^{\min}$ in (9) using the following corollary that follows directly from Theorem 1.

Corollary 1. As a direct result of Theorem 1, we can reduce (10) to

$$\Pr\left\{\beta_i(t) < \mu_k(d+1) - \sqrt{Q_{d+1}(\gamma)\mathbf{e}_{d+1}^T \Sigma_k \mathbf{e}_{d+1}}\right\} < \frac{1-\gamma}{2}. \quad (12)$$

Therefore, we find $D_i^{\min}(\epsilon)$ and ϵ in (9) as

$$D_i^{\min}(\epsilon) = \mu_k(d+1) - \sqrt{Q_{d+1}(1-2\epsilon)\mathbf{e}_{d+1}^T \Sigma_k \mathbf{e}_{d+1}}. \quad (13)$$

Since $Q_{d+1}(\gamma)$ can only be calculated numerically, a closed-form relationship between $D_i^{\min}(\epsilon)$ and ϵ cannot be found. However, we can numerically analyze this relationship, as shown in Fig. 4. Fig. 4 is found using a set of points generated with (13) for different values of $\mu_k(d+1)$ and Σ_k . From Fig. 4, we can first observe that $D_i^{\min}(\epsilon)$ is an increasing function. *This means that the probability of the human brain noticing QoS differences for low delays will be much smaller than for higher delays*, which is an intuitive fact. Furthermore, it can be inferred that, if the delay perception for a group of human users within a cluster is diverse, then the system's confidence on the delay perception of this group of humans will decrease, i.e., the estimation of the delay perception of this group of human users will be less reliable. Next, we determine constraint (2b) using $D_i^{\min}(\epsilon)$.

As stated before, some delays are not perceptible to human users. To capture this feature, we find $D_i^{\max}(\beta_i(t))$ and $\epsilon(\beta_i(t))$ in problem (2) using $D_i^{\min}(\epsilon)$. Recall that $D_i^{\max}(\beta_i(t))$ is a parameter that will be used by the resource allocation system to represent the maximum tolerable delay for the reliable communication of user i with $1 - \epsilon(\beta_i(t))$ being the reliability of user i . There are three possible cases for $D_i^{\max}$ based on $D_i^{\min}$ of a human user i :

1) $D_i^{\max} > D_i^{\min}$: In this case, the system will not be reliable even if we satisfy $\Pr(D > D_i^{\max}) < \epsilon$. The reason is that the human user has a delay perception of less than the maximum delay $D_i^{\max}$ and hence, the system is not reliable.

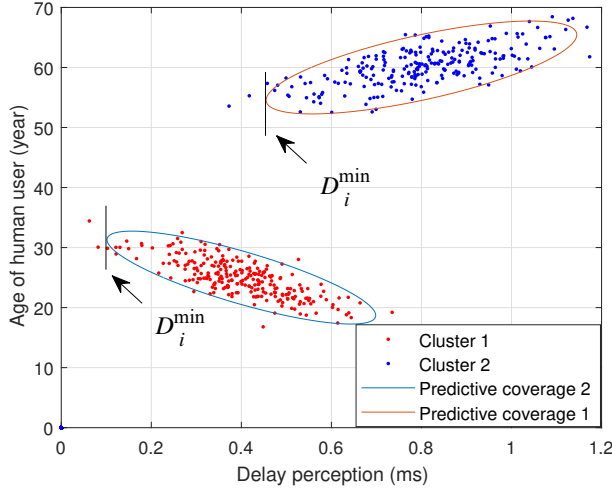


Figure 3: Finding $D_i^{\min}$ using a GMM model for two different clusters.

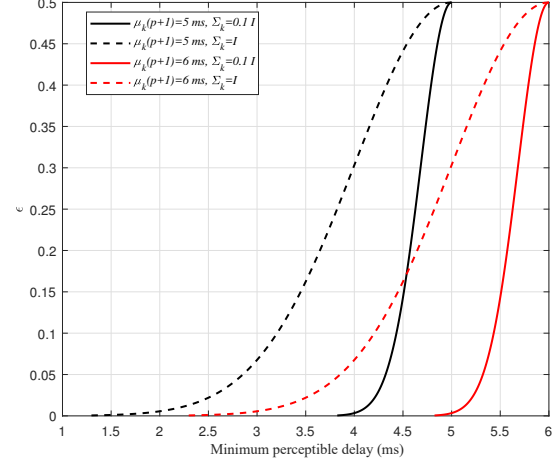


Figure 4: Relationship between ϵ and $D_i^{\min}(\epsilon)$ for different values of $\mu_k(d+1)$ and Σ_k . I is the identity matrix.

2) $D_i^{\max} < D_i^{\min}$: In this case, if the system is able to satisfy $\Pr(D > D_i^{\max}) < \epsilon$, then the system will be reliable, because user i cannot sense delays less than $D_i^{\min}$ and its service delay will not exceed $D_i^{\max}$.

3) $D_i^{\max} = D_i^{\min}$: If this equality holds, the system will be reliable and it will also have prevented a waste of resources. If any given user cannot perceive delays less than $D_i^{\min}$, then it is not effective to allocate more resources to this user.

We define $\mathcal{S}$ as the event resulting from case 2 and case 3 while assuming events E_1 and E_2 satisfy $D < D_i^{\max}$ and $\beta_i(t) > D_i^{\min}$, respectively. We know that for case 1, event $E_1 \cap E_2$ is a subset of event $\mathcal{S}$, and in case 2, event $\mathcal{S}$ is a subset of event $E_1 \cap E_2$. Similarly, in case 3, event $E_1 \cap E_2$ is same as event $\mathcal{S}$. Since the probability of $E_1 \cap E_2$ can be computed, if we set $D_i^{\min}$ to $D_i^{\max}$ (case 3), we can find $\mathcal{S}$ as follows:

$$\Pr(E_1 \cap E_2) = 1 - \Pr\left((D > D_i^{\max}) \cup (\beta_i(t) < D_i^{\min})\right) \quad (14)$$

$$= 1 - \left(\Pr(D > D_i^{\max}) + \Pr(\beta_i(t) < D_i^{\min}) - \Pr(D > D_i^{\max})\Pr(\beta_i(t) < D_i^{\min})\right). \quad (15)$$

(14) follows from De Morgan's law, and (15) is true since D and $\beta_i(t)$ are two independent random variables at each iteration. Therefore, if $D_i^{\min} = D_i^{\max}$ for user i and ϵ, ϵ' is small, we can see that

$$1 - \left(\Pr(D > D_i^{\max}) + \Pr(\beta_i(t) < D_i^{\min})\right) \geq 1 - (\epsilon + \epsilon'), \quad (16)$$

and, hence, $\Pr(\mathcal{S}) = \Pr(E_1 \cap E_2) > 1 - (\epsilon + \epsilon')$, where $\Pr(\mathcal{S})$ is the reliability of the system defined

in (2). Subsequently, as we design the system, we consider the reliability as a predetermined target design parameter for the system. Using this parameter, we can set ϵ and ϵ' . Given ϵ' and numerical function $D_i^{\min}(\epsilon')$ derived in (12), $D_i^{\min}$ can be determined. Now, given $\epsilon(\beta_i(t))$ and $D_i^{\max}(\beta_i(t))$, we can fully characterize problem (2). We note that, as the human user delay perception changes throughout a day, $D_i^{\max}(\beta_i(t))$ changes accordingly. Therefore, our solution should take into account changes in $D_i^{\max}(\beta_i(t))$ as well as changes in channel gains $h_{ij}(t)$. $D_i^{\max}$ is a threshold and $1 - \epsilon$ denotes the reliability which is used in all cellular systems. However, in a wireless system that explicitly takes into account human users in its loop, $D_i^{\max}$ will become a function of $\beta_i(t)$ and can be determined using the concept of effective delay $D_i^{\min}(\epsilon')$ defined in (9).

III. BRAIN-AWARE RESOURCE MANAGEMENT

The notion of human-in-the-loop implies that human factors (such as brain limitations) will be part of the resource allocation framework, i.e., in the loop of resource allocation. For such a system, resource management can dynamically adapt to the human user in its loop, as opposed to just the device. Therefore, our approach considers the human brain limitations as functions of time and adapts the system to the dynamic changes that can occur in the brain and its cognitive limitations, over time.

Since the human brain state usually changes rapidly, our work is different from context-aware or QoS-aware works. The fast fluctuations in the cognitive activities of a human brain which have been validated in many works such as [39]–[41] requires the resource allocation framework to be aware of time-varying brain-aware delay constraint in each time slot.

To solve problem (2), we propose a novel brain-aware resource management framework that takes into account the time-varying wireless channel and the time-varying brain-aware delay constraint (2b). We transform this constraint into a mathematically tractable form.

The relation between the packet length distribution and the service time distribution for a packet is shown next, in Corollary 2. Here, the packet service time is defined as the transmission time of a packet from the BS to the UE or MTD.

Corollary 2. If a fixed rate r_i is allocated to a user and the packet lengths follow an exponential distribution with parameter χ , then, the distribution of the service time s will also be exponential with parameter χr_i .

Proof: The CDF of the exponential distribution is $\mathcal{F}_l(\psi) = \Pr(l < \psi) = 1 - e^{-\chi\psi}$. Hence,

$$\mathcal{F}_s(S) = \Pr(s < S) = \Pr\left(\frac{l}{r_i} < S\right) = \Pr(l < r_i S) = \mathcal{F}_{r_i S}(s) = 1 - e^{-\chi r_i S}. \quad (17)$$

This means that the PDF for the service time is $f_S(s) = e^{-\chi r_i s}$. ■

the packet length distribution parameter χ is constant in our analysis, without loss of generality, we assume that the service time of each packet is an exponential random variable with parameter r_i , which is the same as the rate allocated to the user. We assume that for any given user, the packets arrive according to a Poisson process with the rate $a_i(\tau)$, and the user data rate is exponential with parameter $r_i(\tau)$ in slot $\tau = 1, \dots, t$. Next, we derive the probability with which the delay of a given user i exceeds a threshold $D_i^{\max}$.

Theorem 2. Assume that user i has a time varying rate r_i^τ at time slot τ . If the duration of each time slot is long enough for the queue to reach its steady state, i.e.,

$$\frac{1}{r_i(\tau) - a_i(\tau)} \ll \delta\tau, \quad (18)$$

then, the probability that the delay exceeds a threshold is

$$\Pr(D > D_i^{\max}) = \lim_{t \rightarrow \infty} \frac{1}{t} \sum_{\tau=1}^t e^{-(r_i(\tau) - a_i(\tau)) D_i^{\max}}, \quad (19)$$

under the condition that $r_i(\tau) > a_i(\tau)$ for all $\tau > 0$.

Proof: See Appendix B. ■

Theorem 2 shows that constraint (2b) is satisfied if the network can satisfy the following condition

$$\lim_{t \rightarrow \infty} \frac{1}{t} \sum_{\tau=1}^t e^{-(r_i(\tau) - a_i(\tau)) D_i^{\max}} < \epsilon. \quad (20)$$

Fig. 5 shows the relationship between the theoretical result from Theorem 2 and simulation results. Clearly, simulation and analytical results are a near-perfect match with a maximum error of only 0.0146.

A. Optimal Resource Allocation with Guaranteed Reliability

Constraint (20) is analogous to the drift-plus-penalty method in Lyapunov optimization framework [42] which we use to solve (2). The problem has a time-varying nature since the human brain conditions and needs will change from time to time. The users' processing state $\beta(t)$ is also a function of time, and accordingly, the latency needs in (2b) will be time-varying. Therefore, we need to solve the optimization problem (2) during each time slot efficiently. We propose an

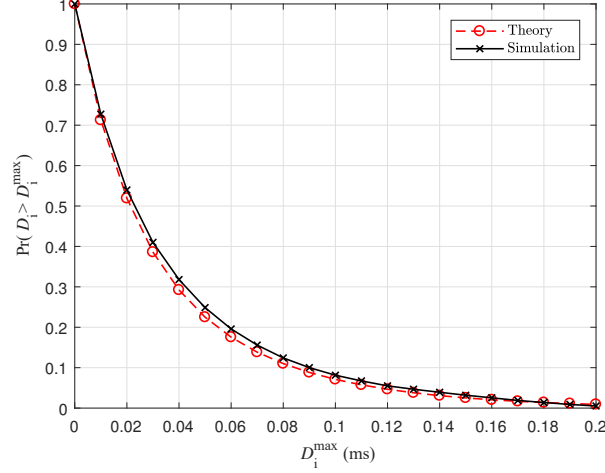


Figure 5: Comparison between simulation results and the result of Theorem 2.

algorithm with a low computational complexity for solving problem (2). The drift-plus-penalty approach is used to stabilize a queue network while minimizing time average of a penalty function. To satisfy constraint (2b) in all time slots, (19) must be smaller than ϵ . For this reason, we use virtual queues to model the time average constraint (20) in the optimization problem. We define a virtual queue:

$$F_i(t+1) = \max\{F_i(t) + e^{-(r_i(t)-a_i(t))D_i^{\max}} - \epsilon, 0\}. \quad (21)$$

We can see that $e^{-(r_i(t)-a_i(t))D_i^{\max}} - \epsilon < F_i(t) - F_i(t)$. Consequently, we obtain:

$$\sum_{\tau=1}^t e^{-(r_i(\tau)-a_i(\tau))D_i^{\max}} - \epsilon t < F_i(t) - F_i(0). \quad (22)$$

If $F_i(0)$ is bounded, we have:

$$\lim_{t \rightarrow \infty} \frac{1}{t} \sum_{\tau=1}^t e^{-(r_i(\tau)-a_i(\tau))D_i^{\max}} - \epsilon < \lim_{t \rightarrow \infty} \frac{F_i(t)}{t}. \quad (23)$$

If the queue $F_i(t)$ is mean-rate stable, that is, $\lim_{t \rightarrow \infty} \frac{F_i(t)}{t} = 0$, then we have:

$$\lim_{t \rightarrow \infty} \frac{1}{t} \sum_{\tau=1}^t e^{-(r_i(\tau)-a_i(\tau))D_i^{\max}} < \epsilon. \quad (24)$$

The Lyapunov function is defined for all the queues in the base station as $Y(t) = \frac{1}{2} \sum_{i \in \text{MUH}} F_i(t)^2$.

Then, we can find the drift function $\Delta t = Y(t+1) - Y(t)$ as:

$$Y(t+1) = \frac{1}{2} \sum_{i \in \text{MUH}} F_i(t+1)^2 \leq \frac{1}{2} \sum_{i \in \text{MUH}} F_i(t)^2 + \frac{1}{2} \sum_{i \in \text{MUH}} y_i(t)^2 + \sum_{i \in \text{MUH}} y_i(t)F_i(t), \quad (25)$$

where

$$y_i^c(t) = e^{-(r_i(t)-a_i(t))D_i^{\max}} - \epsilon. \quad (26)$$

Thus,

$$\Delta t \leq \frac{1}{2} \sum_{i \in \mathcal{M} \cup \mathcal{H}} y_i^c(t)^2 + \sum_{i \in \mathcal{M} \cup \mathcal{H}} y_i^c(t) F_i(t). \quad (27)$$

We can form the drift-plus-penalty by adding $V \sum_{i,j} p_{ij}(t)$ to both sides of inequality (27), where $\sum_{i,j} p_{ij}$ is the total power of the BS which we want to minimize, and V is a parameter that determines how important minimizing the objective function (2a) is in comparison with satisfying (2b). We can balance the tradeoff between power and delay. The drift-plus-penalty inequality is

$$\Delta t + V \sum_{i,j} p_{ij} \leq \frac{1}{2} \sum_{i \in \mathcal{M} \cup \mathcal{H}} y_i^c(t)^2 + V \sum_{i,j} p_{ij}(t) + \sum_{i \in \mathcal{M} \cup \mathcal{H}} y_i^c(t) F_i(t). \quad (28)$$

Given that we assumed $r_i(t) > a_i(t)$ for all t , we know that $|y_i^c(t)| < 1 \forall t, i \in \mathcal{M} \cup \mathcal{H}$, and hence, we can rewrite (28) as

$$\Delta t + V \sum_{i,j} p_{ij} \leq UB + V \sum_{i,j} p_{ij}(t) + \sum_{i \in \mathcal{M} \cup \mathcal{H}} y_i^c(t) F_i(t), \quad (29)$$

where UB is the upper bound of $\frac{1}{2} \sum_{i \in \mathcal{M} \cup \mathcal{H}} y_i^c(t)^2$, and is equal to $\frac{|\mathcal{H}|+|\mathcal{M}|}{2}$. $|\mathcal{H}|$ is the cardinality of set $\mathcal{H}$. Using the drift-plus-penalty algorithm [43], we know that, by minimizing the right hand side of equation (29), queue $F_i(t)$ will be mean-rate stable, and hence, the condition $y_i^c(t) < 0$ will be satisfied. As a result, constraint (2b) will also be satisfied. Furthermore, we know that by minimizing the right hand side of (29), cost function (2a) is also minimized, owing to the fact that (2a) is defined as a penalty function. By minimizing the right hand side of (29), our optimization problem can be converted to the following time-varying problem:

$$\min_{\rho(t), \mathbf{P}(t)} \quad V \sum_{i,j} p_{ij}(t) + \sum_{i \in \mathcal{M} \cup \mathcal{H}} y_i^c(t) F_i(t), \quad (30a)$$

$$\text{s.t.} \quad r_i(t) > a_i(t), \quad \forall i \in \mathcal{H} \cup \mathcal{M} \quad (30b)$$

$$p_{ij}(t) \geq 0, \quad \rho_{ij}(t) \in \{0, 1\}, \quad \forall i \in \mathcal{H} \cup \mathcal{M}, j \in \mathcal{K}, \quad (30c)$$

$$\sum_{i \in \mathcal{H} \cup \mathcal{M}} \rho_{ij}(t) = 1, \quad \forall j \in \mathcal{K}. \quad (30d)$$

The cost function in (30a) is equivalent to (2a) and (2b) in the original optimization problem. Learning the effective delay of each human user using our proposed PDI method determines the parameters $y_i^c(t)$ and $F_i(t)$ in the problem (30a). However, in order to satisfy (2b), we need to also satisfy (30b). The reason for adding (30b) is that if this constraint is not satisfied in

any time slot, the queue length will approach infinity. Constraints (30c) and (30d) are feasibility conditions and remain the same as (2). Hence, by solving (30) in each time slot, the original problem (2) will be solved.

Nonetheless, problem (2a) is not a convex optimization problem, due to the fact that it is a mixed integer problem and its complexity increases exponentially with the number of users. Since (2a) needs to be solved at each time slot, this exponential order of complexity makes the implementation infeasible. Consequently, we should use a dual decomposition method to break down optimization problem (30) to smaller subproblems, and find the optimal solution to (30) using a low complexity method. It is rather challenging to solve (30) using a dual decomposition method, as the structure of $y_i^c(t)$ makes it infeasible to decompose the objective function for each RB. In order to overcome this challenge, we convert (30) to a decomposable form. Then, we will show that this converted problem is equivalent to (30).

For this purpose, the Lagrangian for problem (30) is written as

$$V \sum_{i,j} p_{ij}(t) + \sum_{i \in \mathcal{M} \cup \mathcal{H}} y_i^c(t) F_i(t) + \sum_{i \in \mathcal{M} \cup \mathcal{H}} \lambda_i (a_i(t) - r_i(t)), \quad (31)$$

where λ_i is the Lagrange multiplier. As we know, $y_i^c(t) = e^{-(r_i(t) - a_i(t)) D_i^{\max}} - \epsilon$. Therefore, the only decision variables are allocation of resource blocks to the users and allocating power to each RB. Although $F_i(t)$ is a function of $y_i^c(t)$, it is not a decision variable and is treated as a constant. Hence, (31) can be rewritten as

$$V \sum_{i,j} p_{ij}(t) + \sum_{i \in \mathcal{M} \cup \mathcal{H}} e^{-(r_i(t) - a_i(t)) D_i^{\max}} F_i(t) + \sum_{i \in \mathcal{M} \cup \mathcal{H}} \lambda_i (a_i(t) - r_i(t)). \quad (32)$$

The main optimization problem consists of two components. First, minimizing the total power of the BS with weight V , and second, minimizing the summation $\sum_{i \in \mathcal{M} \cup \mathcal{H}} e^{-r_i(t)}$ which has a weight $F_i(t) e^{-a_i(t) D_i^{\max}}$ for each user i .

As we can see, (32) is not decomposable for each RB. Here we will have an approximation of (30) and then propose an algorithm to solve this approximation efficiently. In this C-additive approximation, $\sum_{i \in \mathcal{M} \cup \mathcal{H}} e^{-(r_i(t) - a_i(t)) D_i^{\max}} F_i(t)$ in (32) is substituted with its linear approximation of exponential term e^{-x} at $x = 0$.

$$\sum_{i \in \mathcal{M} \cup \mathcal{H}} -(r_i(t) - a_i(t)) D_i^{\max} F_i(t). \quad (33)$$

In the original problem, if $y_i^c(t)$ starts to become greater than zero for user i , then $F_i(t)$ will increase and it will give more weight to the term $e^{-(r_i(t) - a_i(t)) D_i^{\max}}$. As a result, the algorithm

allocates more resources to user i such that it minimizes $e^{-(r_i(t)-a_i(t))D_i^{\max}}$ for user i , and accordingly, $y_i^c(t)$ decreases. Hence, $F_i(t)e^{-(r_i(t)-a_i(t))D_i^{\max}}$ plays the role of feedback in the system. As we can see from (33), this approximation will not change this feedback mechanism and plays the same role in the system. Therefore, we can write

$$\begin{aligned} & \min_{\mathbf{P}, \rho} \left\{ V \sum_{i,j} p_{ij}(t) + \sum_{i \in \mathcal{M} \cup \mathcal{H}} -(r_i(t) - a_i(t)) D_i^{\max} F_i(t) \right\} \\ & < C + \min_{\mathbf{P}, \rho} \left\{ V \sum_{i,j} p_{ij}(t) + \sum_{i \in \mathcal{M} \cup \mathcal{H}} e^{-(r_i(t)-a_i(t))D_i^{\max}} F_i(t) \right\}. \end{aligned} \quad (34)$$

Using this C-additive approximation, it can be easily proved that all terms are mean-rate stable. Hence, (2b) in the original problem is satisfied [42]. Finally, problem (2) can be presented as:

$$\begin{aligned} & \min_{\rho(t), \mathbf{P}(t)} \quad V \sum_{i,j} p_{ij}(t) - \sum_{i \in \mathcal{M} \cup \mathcal{H}} (r_i(t) - a_i(t)) D_i^{\max} F_i(t), \\ & \text{s.t.} \quad r_i(t) > a_i(t), \end{aligned} \quad (35a)$$

$$p_{ij}(t) \geq 0, \quad \forall i \in \mathcal{H} \cup \mathcal{M}, j \in \mathcal{K}, \quad (35b)$$

$$\rho_{ij}(t) \in \{0, 1\}, \quad \forall i \in \mathcal{H} \cup \mathcal{M}, j \in \mathcal{K}, \quad (35c)$$

$$\sum_{i \in \mathcal{H} \cup \mathcal{M}} \rho_{ij}(t) = 1, \quad \forall j \in \mathcal{K}. \quad (35d)$$

In order to solve this problem, we can decompose it into K subproblems. Since these subproblems are coupled through constraint (35d), we use the dual decomposition method for solving (35) [44]. First, the Lagrangian is written for problem (35), and in the second step, it is decomposed for each RB. Then, the resource block allocation and the power of each RB are found in terms of the Lagrange multiplier vector $\boldsymbol{\lambda}$. Finally, $\boldsymbol{\lambda}$ is calculated using an ellipsoid method.

The Lagrangian for problem (35) is

$$\begin{aligned} \mathcal{L}_a(\mathbf{P}, \rho, \boldsymbol{\lambda}) &= V \sum_{i,j} p_{i,j}(t) + \sum_{i \in \mathcal{M} \cup \mathcal{H}} -(r_i(t) - a_i(t)) D_i^{\max} F_i(t) - \lambda_i (r_i(t) - a_i(t)) \\ &= V \sum_{i,j} p_{i,j}(t) - \sum_{i \in \mathcal{M} \cup \mathcal{H}} (\lambda_i + D_i^{\max} F_i(t)) (r_i(t) - a_i(t)). \end{aligned} \quad (36)$$

One major difference between our problem and conventional power minimization problems is that there is an additional term $D_i^{\max} F_i(t)$ added to the Lagrange multiplier (the shadow price).

In this problem, $D_i^{\max} F_i(t)$ plays the role of a bias term. Therefore, a new hypothetical Lagrange multiplier λ'_i is assumed and defined as $\lambda'_i = \lambda_i + D_i^{\max} F_i(t)$. This means that adding constraint (2b) to the problem instead of constraint (35a) increases the shadow price by a factor of $D_i^{\max} F_i(t)$. Increasing the shadow price for a constraint makes it looser. As a result, in many

Algorithm 3 Resource allocation algorithm

- 1: Obtain $D_i^{\max(t)}$, $\epsilon_i(t) \forall i \in \mathcal{H} \cup \mathcal{M}$ using PDI algorithm (Algorithm 2).
 - 2: Find $F_i(t)$, $\forall i \in \mathcal{H} \cup \mathcal{M}$, using (21)
 - 3: Initialize λ
 - 4: **while** convergence condition is not satisfied **do**
 - 5: Find p_{ij} , $\forall i \in \mathcal{H} \cup \mathcal{M}, j \in \mathcal{K}$, using the updated λ (40)
 - 6: for each RB $j \in \mathcal{K}$, find $k(j)$ by searching over all users $i \in \mathcal{H} \cup \mathcal{M}$ using (41) and then assign ρ_{ij}^* and p_{ij}^* for all $i \in \mathcal{H} \cup \mathcal{M}, j \in \mathcal{K}$,
 - 7: Use the ellipsoid method to find λ
 - 8: **end while**
-

time slots, constraint (35a) will not be a tight constraint and the Lagrange multiplier will be set to $\lambda_i = [\lambda'_i - D_i^{\max} F_i(t)]^+$. the Lagrange dual function is

$$g(\lambda) = \min_{\rho(t), \mathbf{P}(t)} \mathcal{L}_a(\mathbf{P}, \rho, \lambda). \quad (37)$$

The minimization problem (37) can be decomposed to K subproblems. $g'_j(\lambda)$ can be written as

$$g'_j(\lambda) = \min_{\mathbf{P}(t)} V \sum_i p_{i,j} - \sum_{i \in \mathcal{M} \cup \mathcal{H}} (\lambda_i + D_i^{\max} F_i(t)) (W \log_2(1 + K \frac{h_{i,j} p_{i,j}}{\sigma^2})), \quad (38)$$

where $\mathcal{D}$ is a set of feasible p_{ij} s in which for RB j , there is only one i that $p_{ij} \neq 0$. Hence, $g(\lambda)$ is

$$g(\lambda) = \sum_j g'_j(\lambda) + \sum_{i \in \mathcal{M} \cup \mathcal{H}} (\lambda_i + D_i^{\max} F_i(t)) (a_i(t)). \quad (39)$$

If λ is fixed, $g'_j(\lambda)$ is a convex function of $\mathbf{P}$. Therefore, $\mathbf{P}$ is found by taking a derivate with respect to p_{ij} and setting it to zero. This results in

$$p_{ij} = \left[\frac{(\lambda_i + D_i^{\max} F_i(t)) W}{V \log_2} - \frac{\sigma^2}{K h_{ij}} \right]^+. \quad (40)$$

The optimal RB allocation for RB j is $k(j)$, and can be written as

$$k(j) = \underset{i}{\operatorname{argmin}} V \sum_i p_{i,j} - \sum_{i \in \mathcal{M} \cup \mathcal{H}} (\lambda_i + D_i^{\max} F_i(t)) (W \log_2(1 + K \frac{h_{i,j} p_{i,j}}{\sigma^2})), \quad (41)$$

$$g'_j(\lambda) = \min_i V \sum_i p_{i,j} - \sum_{i \in \mathcal{M} \cup \mathcal{H}} (\lambda_i + D_i^{\max} F_i(t)) (W \log_2(1 + K \frac{h_{i,j} p_{i,j}}{\sigma^2})). \quad (42)$$

Thus, ρ_{ij}^* and p_{ij}^* will be given by:

$$\rho_{ij}^* = \begin{cases} 1, & i = k(j), \\ 0, & \text{otherwise.} \end{cases} \quad p_{ij}^* = \begin{cases} p_{ij}, & i = k(j), \\ 0, & \text{otherwise.} \end{cases} \quad (43)$$

Hence, the optimal rate becomes $r_i^* = \sum_j W \log_2(1 + K \frac{h_{i,j} p_{i,j}^*}{\sigma^2})$. The only parameter that affects this joint RB and power allocation is λ . As the number of RBs increases, the duality gap in this problem approaches zero [44]. We know that the optimal value is found by maximization

of $g(\boldsymbol{\lambda})$ with respect to $\boldsymbol{\lambda}$. In order to find $\boldsymbol{\lambda}$, we use the ellipsoid method [45], and to do so, we have to find the sub-gradient for the dual objective $g(\boldsymbol{\lambda})$. The following theorem will show that the subgradient for (36) is a vector with elements $\zeta_i = a_i - r_i^*$.

Theorem 3. The subgradient of the dual optimization problem with dual objective defined in (39), is the vector $\mathbf{d}$ whose elements $\zeta_i, \forall i \in \mathcal{H} \cup \mathcal{M}$ are given by:

$$\zeta_i = \begin{cases} a_i - r_i^*, & a_i \geq r_i^*, \\ 0, & a_i < r_i^*. \end{cases} \quad (44)$$

Proof: Since

$$g(\boldsymbol{\lambda}) = \min_{\mathbf{P}, \boldsymbol{\rho}} \mathcal{L}_a(\mathbf{P}, \boldsymbol{\rho}, \boldsymbol{\lambda}) = \mathcal{L}_a(\mathbf{P}^*, \boldsymbol{\rho}^*, \boldsymbol{\lambda}), \quad (45)$$

we have:

$$\begin{aligned} g(\boldsymbol{\delta}) &\leq \mathcal{L}_a(\mathbf{P}^*, \boldsymbol{\rho}^*, \boldsymbol{\delta}) \\ &= V \sum_{i,j} p_{i,j}^*(t) - \sum_{i \in \mathcal{M} \cup \mathcal{H}} (\delta_i + D_i^{\max} F_i(t)) (r_i^*(t) - a_i(t)) \\ &= V \sum_{i,j} p_{i,j}^*(t) - \sum_{i \in \mathcal{M} \cup \mathcal{H}} (\lambda_i + D_i^{\max} F_i(t)) (r_i^*(t) - a_i(t)) \\ &\quad + (\lambda_i - \delta_i) (r_i^*(t) - a_i(t)) = g(\boldsymbol{\lambda}) + (\boldsymbol{\lambda} - \boldsymbol{\delta})^T \boldsymbol{\zeta}', \end{aligned} \quad (46)$$

where $\boldsymbol{\zeta}' = [r_1^* - a_1 \quad \cdots \quad r_{N+M}^* - a_{N+M}]^T$.

However, because of the term $D_i^{\max} F_i(t)$, when $\lambda_i = 0$ and $a_i < r_i^*$, the direction of $\boldsymbol{\zeta}'$ will be infeasible. Using the projected subgradient method [46], we can transform this infeasible direction to a feasible one. The update rule for projected subgradient is: $\boldsymbol{\lambda}^{(k+1)} = \Pi(\boldsymbol{\lambda}^{(k)} - \alpha_k \boldsymbol{\zeta}'_k)$ where α_k is the step size and Π is the Euclidan projection on the feasible set. Since the feasible set is $\lambda_i > 0$, we can see that

$$\Pi(\boldsymbol{\lambda}^{(k)} - \alpha_k \boldsymbol{\zeta}'_k) = \boldsymbol{\lambda}^{(k)} - \alpha_k \boldsymbol{\zeta}_k, \quad (47)$$

where:

$$\zeta_i = \begin{cases} \zeta'_i, & \zeta'_i \geq 0 \\ 0, & \zeta'_i < 0 \end{cases} = \begin{cases} a_i - r_i^*, & a_i \geq r_i^*, \\ 0, & a_i < r_i^*. \end{cases} \quad (48)$$

■

Algorithm 3 summarizes our proposed resource allocation algorithm.

B. Complexity Analysis

Next, we find the complexity of our algorithm which needs to be run in each iteration. There are K RBs in our problem, for each of which (41) needs to be evaluated for $M + N$ users. It takes $\mathcal{O}((M + N)K)$ times to solve a primal problem. Subsequently, the dual problem will be solved, which gives us the optimal value of λ in an $M + N$ dimensional space and has a complexity of $\mathcal{O}((M + N)^2)$. Therefore, the overall complexity should be $\mathcal{O}((M + N)^3K)$. However, as mentioned before, adding $D_i^{\max}F_i(t)$ to the Lagrange multiplier sets a major part of it to zero, and as a result, the order of complexity will decrease to $\mathcal{O}((M + N)K)$. Given the low-complexity of the proposed algorithm, in practice, it can be easily run periodically by the network at each time slot t , so as to effectively adapt to dynamic, time-varying changes in both the human delay perceptions of the users and the wireless channel.

IV. SIMULATION RESULTS AND ANALYSIS

For our simulations, we consider the dataset in [33] to model the delay perception of a human user. In [33], the authors conducted human subject studies using 30 human users, where each subject is asked to rate the quality of 5 movies while the delay and packet loss in the system is being increased. We used the average score of each human user to estimate their delay perception. In [33], The highest delay that the test subject is not able to sense is considered as the delay perception of this subject, and, this, matches our definition of delay perception. We also used a variation of the bootstrap method [38] to increase the number of data points to 1000. We can see the histogram of the delay perception for these 1000 data points in Fig. 7.

To the best of our knowledge, no dataset which includes features for each human user as well as human delay perception currently exists. Hence, we attribute three continuous features to each user. The process of adding features starts by clustering the delay perceptions $\beta_i(t)$ for 1000 users. Then, we choose a random mean vector and a random positive semidefinite covariance matrix for each cluster and use them to create multivariate random Gaussian features for each data in the cluster. In consequence the random features have: 1) a GMM structure and 2) a predictive ability for $\beta_i(t)$. Hence, each user is associated with a vector $\mathbf{w} \in \mathbb{R}^4$.

We consider a network with a bandwidth of 10 MHz, $a_i(t) = 1$ Mbps, $\sigma^2 = -173.9$ dBm, and $\epsilon = 0.05$. We use a circular cell with the cell radius of 1.5 km. We set the path loss exponent to 3 (urban area) and the carrier frequency to 900 MHz. The packet length is an exponential random variable with an average size of 10 kbits. We use 5 MTD and 5 UE in the system

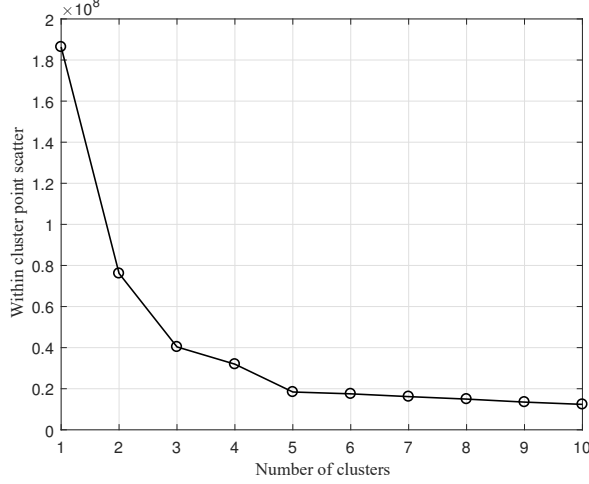


Figure 6: Within point scatter for the EM clustering method on the dataset.

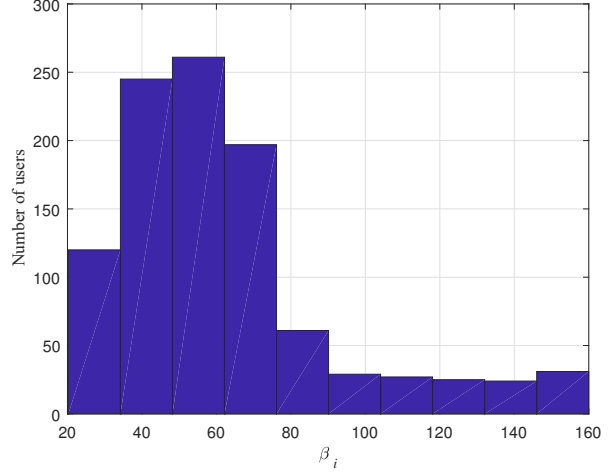


Figure 7: Distribution of $\beta_i(t)$ for the 1000 users in dataset.

and we set $D_i^{\max}$ to 20 ms for them, unless otherwise mentioned. For the brain aware users, we arbitrarily select 5 UE in the system out of all data points. The brain-unaware case is QoS-aware and power-aware. In this case, $D_i^{\max}$ and ϵ are not functions of $\beta_i(t)$.

Fig. 6 shows the within cluster point scatter for the EM algorithm in our dataset. This *within cluster point scatter* for a clustering C is defined as [38]: $W(C) = \frac{1}{2} \sum_{k=1}^n \sum_{c(i)=k} \sum_{c(i')=k} d(x_i, x_{i'})$, where d is an arbitrary distance metric. In essence, the within cluster point scatter is a loss function that allows the determination of hyper-parameters in the clustering algorithm. The hyper-parameter that we seek to find here is the number of clusters in the dataset. As we can see from Fig. 6, after the number of clusters reaches 5, increasing the number of clusters does not decrease the within cluster point scatter substantially. Hence, the optimal number of clusters is 5. This method of model selection known as elbow method allows the algorithm to avoid overfitting.

Fig. 8 shows the total BS power resulting from the proposed brain-aware case and from a brain-unaware case in which UEs have a fixed constraint (2b) with $D_i^{\max}$ between 10 ms to 60 ms. Here, the total power is the objective of main optimization problem (2). Fig. 8 shows that, as the latency increases, the total power decreases, because it is easier to satisfy constraint (2b) at higher latencies. Also, at higher delays, being brain-aware will no longer yield substantial gains, since $\beta_i(t)$ and $D_i^{\max}$ become close to each other and learning $\beta_i(t)$ cannot save resources for the system. In contrast, in Fig. 8, we can see that for stringent low-latency requirements, the proposed brain-aware approach yields significant gains in terms of saving power. In particular, for 10 ms delay in (2b), Fig. 8 shows that the BS in brain-unaware approach uses 44 % more

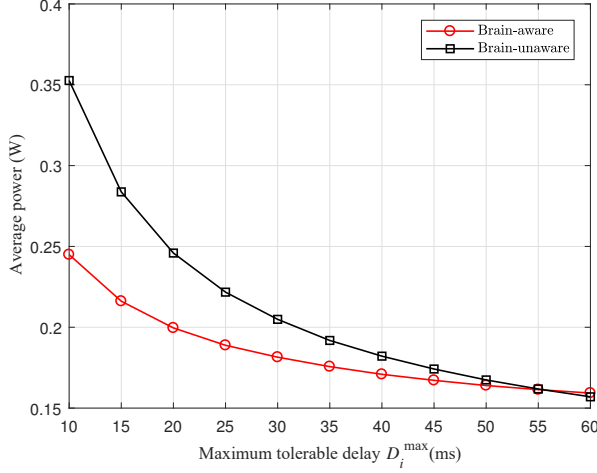


Figure 8: Average power usage of the system as function of different latency requirements for the users.

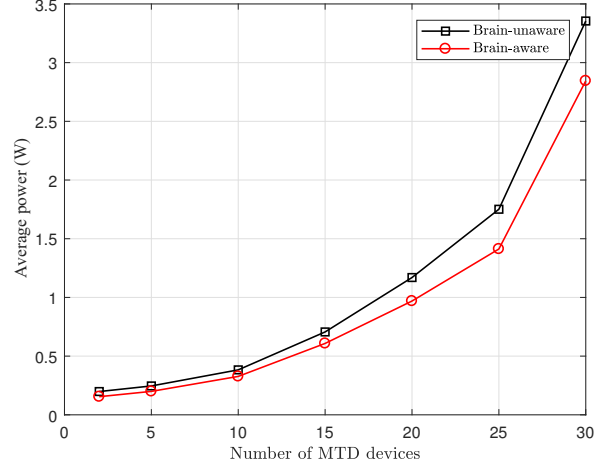


Figure 9: Average power usage of the system for different number of MTDs and 5 UEs with $D_i^{\max} = 20$ ms.

power compared to the brain-aware case. These results stem from the fact that a brain-aware approach can minimize waste of resources and provide service to the users more precisely based on their real brain processing power. Fig. 9 shows average BS power for different number of MTDs. As we can see from Fig. 9, the brain-aware approach will always outperform the brain-unaware approach as the number of MTD increases. For the case of 30 MTD user, the BS in brain-unaware approach uses 16% more power compared to the brain-aware case. This is due to fact that brain-aware approach can allocate resources more efficiently in case of a shortage in resources.

In Fig. 10, we show the average power usage of the system when the number of UEs increases from 2 to 30 with $D_i^{\max}$ set to 20 ms. As the number of users increases, the average power consumption of the system will also increase. This is due to the fact that increasing the number of users will decrease the bandwidth per user. Since the delay and rate requirements of each user are still unchanged, the system needs to use more power to compensate for the bandwidth deficiency. From Fig. 10, we can see that, in the case of 30 users, the brain-aware system is able to save 6.7 dB (78%) on average in the BS power. The brain-aware system can allocate resources based on each user's actual requirement instead of the predefined metrics and this leads to this significant saving in the power consumption of the BS.

In Fig. 11, we show the average power consumed in the system for different number of virtual reality (VR) users. For the VR simulations, we assumed an arrival rate of 25.31 Mbps for each user [47] and have used bandwidths of 20 MHz and 40 MHz. We can see that, the system is

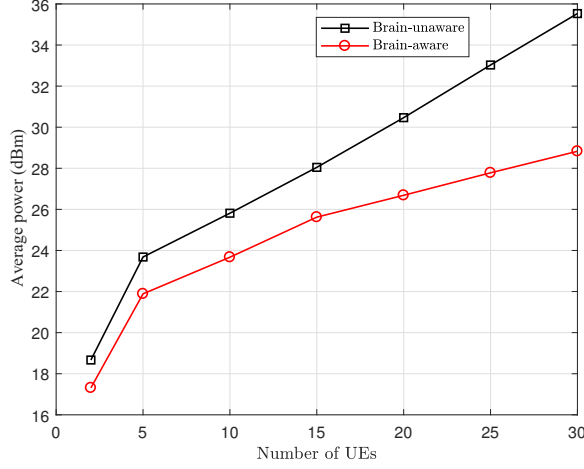


Figure 10: Average power usage of the system for different number of UEs with $D_i^{\max} = 20$ ms.

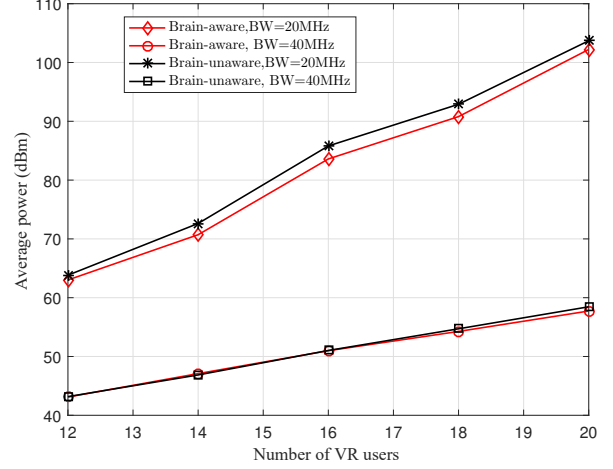


Figure 11: The effect of number of VR users with the rate of 25.31 Mbps and $D_i^{\max} = 20$ ms on the power usage of the system.

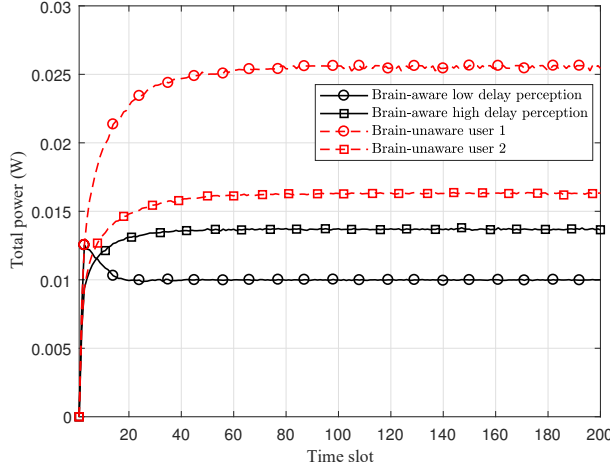


Figure 12: Transmit power for 4 different users. The delay perception of two of the users is learned. Low and high delay perception users have delay perception of 26.8 ms and 133.73 ms, respectively.

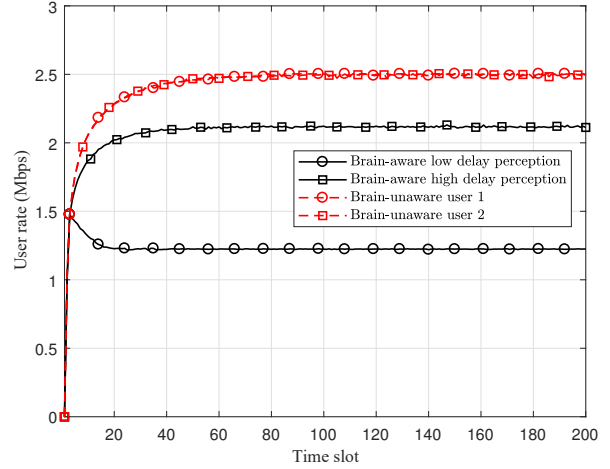


Figure 13: Transmission rate for four different users. The delay perception of two of the users is learned. Low and high delay perception users have delay perceptions of 26.8 ms and 133.73 ms, respectively.

able to save power up to 40% and 15% compared to the brain-unaware scenario in the case of 20 MHz, and 40 MHz bandwidth, respectively. Fig. 11 also shows that the proposed approach is able to allocate resources more efficiently when resources are scarce, i.e. in the 20 MHz case. Also, we can see that increasing the bandwidth will decrease the total power usage in the system which is an inherent feature of communication systems.

In Fig. 12, Fig. 13, and Fig. 14, we consider the case of 7 UEs and 5 MTDs. Two UEs are

chosen as brain-aware users and their delay perception is learned by the PDI method. One of the brain-aware UEs has a delay perception of $\beta_i(t) = 133.73$ ms, and the other one has $\beta_i(t)$ equal to 26.8 ms. The system does not learn the delay perception of the 5 remaining UEs and, hence, it allocates resources to them by using a predefined delay requirement (brain-unaware users).

As we can see in Fig. 12, the power consumption of the first two brain-aware users will be less than that of the brain-unaware users. Moreover, the power consumption for a user with higher delay perception will be less than that of a user with lower delay perception. This shows that the system can successfully allocate resources according to the delay perception of the users. Furthermore, the power consumption related to each user with predetermined delay requirements is different, due to their different channel gains. However, as we will see later, the system is robust to such differences and can guarantee the reliability and rate requirements for users having different channel gains.

In Fig. 13, we show the transmission rate for four different users. We can see that the rate for brain-unaware users with predetermined delay will converge to 2.5 Mbps. This rate will ensure the reliability for these users. However, the rate of the users with learned delay perception will converge to a smaller rate. This is due to the fact that these users' actual requirements are known to the system, and the system uses this knowledge to avoid unnecessarily wasting resources. However, as we will see next, this rate reduction does not change the reliability for these users.

Fig. 14 shows the reliability for the four aforementioned users. As we can see, the reliability of all the users will converge to 95 %, which is the target reliability value for the users. We can see that the system is able to ensure reliability for the users with identified delay perceptions as well as the users with predefined delay requirements. However, as observed from Fig. 12, the system uses 45% less power for those users for which the delay perception is learned.

Finally, Fig. 15 investigates the effect of parameter V for the system with 5 MTDs and 5 UEs. We can see that, as V increases from 1 to 1.9, the convergence time decreases from 40 iterations to 15 iterations. Nevertheless, increasing V will make the algorithm unstable, and as we can see, increasing it to 2.2 will create an overshoot which is 11% higher than the final value. Hence, parameter V , if adjusted correctly, can create a balance between stability and convergence rate of our algorithm.

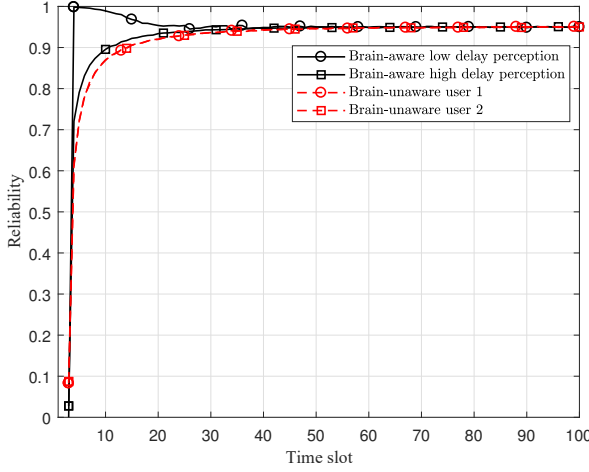


Figure 14: Reliability for 4 different users. The delay perception of two of the users is learned.

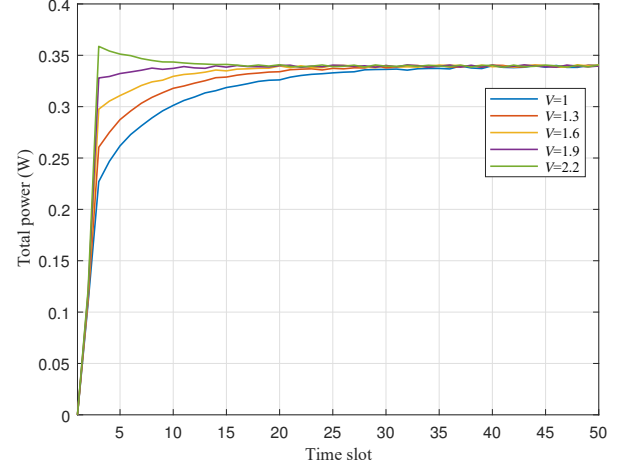


Figure 15: Effect of balancing parameter V in (30a) on the convergence of the resource allocation algorithm.

V. CONCLUSION

In this paper, we have introduced and formulated the notion of delay perception of a human brain, in wireless networks with humans-in-the-loop. Using this notion, we have defined the concept of effective delay of human brain. To quantify this effective delay, we have developed a learning method, named PDI, which consists of an unsupervised and supervised learning part. We have then shown that PDI can predict the effective delay for the human users and find the reliability of this prediction. Then, we have derived a closed-form relationship between the reliability measure and wireless physical layer metrics. Next, using this relationship and the PDI method, we have proposed a novel approach based on Lyapunov optimization for allocating radio resources to human users while considering the reliability of both machine type devices and human users. Our results have shown that the proposed brain-aware approach can save a significant amount of power in the system, particularly for low-latency applications and congested networks. To our best knowledge, this is the first study on the effect of human brain limitations in wireless network design. This paper only scratched the surface of an emerging research area that admits several future extensions. On the one hand, we can extend the studied framework to accommodate other brain-related features beyond the mode of the brain. Examples of such features include perceptual memory and consistency constraints. On the other hand, we can develop recurrent neural network models to capture how the sequence in the brain mode can dynamically change. Finally, another important future work is to conduct real-world experiments with actual users to gather empirical data on brain behavior so as to refine the developed solution.

APPENDIX A

PROOF OF THEOREM 1

We assume that a single brain mode is dominant for each user at each time. We index this single mode as k . For each user i with this dominant mode, $\mathbf{w}_i = [w_1, \dots, w_{d+1}]$ has the following probability density function:

$$p(\mathbf{w}_i) = |2\pi\mathbf{\Sigma}_k|^{-\frac{1}{2}} \exp \left[-\frac{1}{2}(\mathbf{w}_i - \boldsymbol{\mu}_k)^T \mathbf{\Sigma}_k^{-1} (\mathbf{w}_i - \boldsymbol{\mu}_k) \right]. \quad (49)$$

We want to find the smallest region $\mathcal{D}$ in $\mathbb{R}^{d+1}$, in which the delay perception lies with probability γ , i.e.,

$$\int \cdots \int_{\mathcal{D}} p(w_1, w_2, \dots, w_{d+1}) dw_1 \cdots dw_{d+1} = \gamma. \quad (50)$$

$\mathcal{D}$ is not a unique region. However, the objective is to find the smallest region. To this end, we need to find the region where $f(w_1, w_2, \dots, w_{d+1})$ has the greatest value, i.e., if

$$\int \cdots \int_{\mathcal{D}_1} p(w_1, w_2, \dots, w_{d+1}) dw_1 \cdots dw_n = \int \cdots \int_{\mathcal{D}_2} p(y_1, y_2, \dots, y_n) dy_1 \cdots dy_n, \quad (51)$$

and also

$$p(y_1, y_2, \dots, y_{d+1}) \leq p(w_1, w_2, \dots, w_{d+1}) \quad \forall \mathbf{y} \in \mathcal{D}_2, \forall \mathbf{w}_i \in \mathcal{D}_1, \quad (52)$$

then

$$\int \cdots \int_{\mathcal{D}_1} dw_1 \cdots dw_{d+1} \leq \int \cdots \int_{\mathcal{D}_2} dy_1 \cdots dy_{d+1}, \quad (53)$$

which implies that the volume of the region $\mathcal{D}_1$ is smaller than the volume of $\mathcal{D}_2$. Hence, if we find the region $\mathcal{D}$ for which (50) holds, and, using (52), show that all other regions for which (50) holds have greater volumes, then, we would have found the smallest region $\mathcal{D}$, in which the human behavior will stay with the probability γ .

Since $\mathbf{w}_i$ is distributed according to a multivariate Gaussian, we can find the region where it has the highest probability density, i.e., $\{\mathbf{w}_i | p(\mathbf{w}_i) > C_1\}$. This region can be written as:

$$\left\{ \mathbf{w}_i \left| |2\pi\mathbf{\Sigma}_k|^{-\frac{1}{2}} \exp \left[-\frac{1}{2}(\mathbf{w}_i - \boldsymbol{\mu}_k)^T \mathbf{\Sigma}_k^{-1} (\mathbf{w}_i - \boldsymbol{\mu}_k) \right] > C_1 \right. \right\}, \quad (54)$$

which is equivalent to

$$\mathcal{D} = \left\{ \mathbf{w}_i \left| (\mathbf{w}_i - \boldsymbol{\mu}_k)^T \mathbf{\Sigma}_k^{-1} (\mathbf{w}_i - \boldsymbol{\mu}_k) < C_2 \right. \right\}, \quad (55)$$

where C_2 is a positive constant and equals $-\ln |2\pi\mathbf{\Sigma}_k|^{\frac{1}{2}} C_1$. Since $\mathbf{\Sigma}_k$ is a positive definite matrix, (55) is the inner volume of an ellipsoid in a d dimensional space.

We now conjecture that this ellipsoid $\mathcal{D}$ is the smallest region, in which the delay perception

lies with probability γ , i.e., the probability of $\mathbf{w}_i$ being in this region is γ . We use a proof by contradiction to show this. Consider that there exists any other space $\mathcal{E}$ which is smaller than $\mathcal{D}$, and the probability of $\mathbf{w}_i$ being in this region is γ . We can partition $\mathcal{E}$ into two parts $\mathcal{A} = \mathcal{E} \cap \mathcal{D}$ and $\mathcal{E}_2 = \mathcal{E} \cap \mathcal{D}'$, where $\mathcal{D}'$ is the complement of the set $\mathcal{D}$. We also define $\mathcal{D}_2 = \mathcal{D} \cap \mathcal{E}'$. We know that

$$\int_{\mathcal{D}} p(\mathbf{w}_i) d\mathbf{w}_i = \int_{\mathcal{A}} p(\mathbf{w}_i) d\mathbf{w}_i + \int_{\mathcal{D}_2} p(\mathbf{w}_i) d\mathbf{w}_i \quad (56)$$

$$= \int_{\mathcal{E}} p(\mathbf{w}_i) d\mathbf{w}_i = \int_{\mathcal{A}} p(\mathbf{w}_i) d\mathbf{w}_i + \int_{\mathcal{E}_2} p(\mathbf{w}_i) d\mathbf{w}_i = \gamma. \quad (57)$$

Hence, $\int_{\mathcal{D}_2} f(\mathbf{w}_i) d\mathbf{w}_i = \int_{\mathcal{E}_2} f(\mathbf{w}_i) d\mathbf{w}_i$. Since

$$p(\mathbf{w}_i) < C_1 \leq p(\mathbf{y}) \quad \forall \mathbf{w}_i \in \mathcal{E}_2, \mathbf{y} \in \mathcal{D}_2, \quad (58)$$

using (51) and (52) we have $\int_{\mathcal{E}_2} d\mathbf{w}_i < \int_{\mathcal{D}_2} d\mathbf{w}_i$. This means that the set $\mathcal{E}$ has a bigger volume than $\mathcal{D}$, which is a contradiction to our first assumption. This proves that region $\mathcal{D}$ is the smallest region in $\mathbb{R}^{d+1}$ that has the probability γ .

Next, we find the relation between C_2 and γ . γ can be defined as $\int_{\mathcal{D}} p(\mathbf{w}_i) d\mathbf{w}_i$ and can be calculated using chi-square distribution [48]. The region $\mathcal{D}$ can be written as

$$\mathcal{D} = \{\mathbf{w}_i | (\mathbf{w}_i - \boldsymbol{\mu}_k)^T \boldsymbol{\Sigma}_k^{-1} (\mathbf{w}_i - \boldsymbol{\mu}_k) \leq Q_{d+1}(\gamma)\}, \quad (59)$$

where $Q_{d+1}(\gamma)$ is the quantile function of the chi-square distribution with $d + 1$ degrees of freedom. It is defined as $Q_{d+1}(\gamma) = \inf \left\{ x \in \mathbb{R} | \gamma \leq \int_0^x \chi_d^2(u) du \right\}$.

Having defined the confidence region based on γ , we now must find the edges of this ellipsoid. We know that the center of this ellipsoid is $\boldsymbol{\mu}_k$. We need to solve the following optimization problem:

$$\min_{\mathbf{w}_i} \text{ or } \max_{\mathbf{w}_i} \mathbf{e}_j^T \mathbf{w}_i, \quad \text{subject to } \mathbf{w}_i \in \mathcal{D}, \quad (60)$$

where $\mathbf{e}_j$ is a unit vector in $\mathbb{R}^{d+1}$, having 1 in its j th element and zero otherwise. Using KKT conditions for solving the above problem, we have:

$$\mathbf{e}_j + \lambda \boldsymbol{\Sigma}_k^{-1} (\mathbf{w}_i - \boldsymbol{\mu}_k) = 0, \quad (61a)$$

$$(\mathbf{w}_i - \boldsymbol{\mu}_k)^T \boldsymbol{\Sigma}_k^{-1} (\mathbf{w}_i - \boldsymbol{\mu}_k) \leq Q_{d+1}(\gamma), \quad (61b)$$

$$\lambda \left((\mathbf{w}_i - \boldsymbol{\mu}_k)^T \boldsymbol{\Sigma}_k^{-1} (\mathbf{w}_i - \boldsymbol{\mu}_k) - Q_{d+1}(\gamma) \right) = 0, \quad \lambda \geq 0. \quad (61c)$$

The inequality in (61b) is tight. With some algebraic manipulation, we have $\mathbf{w}_i - \boldsymbol{\mu}_k(j) = \frac{1}{\lambda} \boldsymbol{\Sigma} \mathbf{e}_j$,

and so, $\frac{1}{\lambda^2} \mathbf{e}_j^T \Sigma_k \Sigma_k^{-1} \Sigma_k \mathbf{e}_j = Q_{d+1}(\gamma)$. Therefore $\mathbf{w}_i = \pm \sqrt{\frac{Q_{d+1}(\gamma)}{\mathbf{e}_j^T \Sigma_k \mathbf{e}_j}} \Sigma_k \mathbf{e}_j + \boldsymbol{\mu}_k$, $\lambda = \pm \sqrt{\frac{\mathbf{e}_j^T \Sigma_k \mathbf{e}_j}{Q_{d+1}(\gamma)}}$, and $\mathbf{e}_j^T \mathbf{w}_i = \pm \sqrt{Q_{d+1}(\gamma)} \mathbf{e}_j^T \Sigma_k \mathbf{e}_j + \mu_k(j)$.

If λ is positive, we can find the maximum which is $+\sqrt{Q_{d+1}(\gamma)} \mathbf{e}_j^T \Sigma_k \mathbf{e}_j + \mu_k(j)$, and if λ is negative, we can find the minimum which is $-\sqrt{Q_{d+1}(\gamma)} \mathbf{e}_j^T \Sigma_k \mathbf{e}_j + \mu_k(j)$.

Here, $\mu_k(j)$ is the j th element of $\boldsymbol{\mu}_k$. If we set $j = d + 1$, then the delay perception of user j is in the following range:

$$-\sqrt{Q_{d+1}(\gamma)} \mathbf{e}_{d+1}^T \Sigma_k \mathbf{e}_{d+1} < \beta_i(t) - \mu_k(d + 1) < \sqrt{Q_{d+1}(\gamma)} \mathbf{e}_{d+1}^T \Sigma_k \mathbf{e}_{d+1}, \quad (62)$$

at least with probability γ . Hence, Theorem 1 is proved.

APPENDIX B

PROOF OF THEOREM 2

Since the queuing delay is much smaller than the duration of each time slot, we can assume that each packet arriving at a specific time slot will be served at the same time slot. For analyzing the packet delay, we consider a packet that just arrives in the system in time slot τ_k , and find $\Pr(D > D_i^{\max})$ for this packet. When this packet arrives, there are m packets in the system. From lemma 2, we know that the serving time will be an exponential random variable. Since the exponential distribution is memoryless, there is no distinction between a packet already in service and the other packets. Therefore, the waiting time for the packet that has just arrived is the summation of m exponential distributions. Also, the transmission delay for this packet will be another exponential random variable. Hence, the delay of a packet which arrives at time slot τ_k while there are m packets in the system can be written as:

$$d(\tau_k, m) = t_s + t_1(\tau_k) + t_2(\tau_k) + \cdots + t_{m-1}(\tau_k) + t_c(\tau_k), \quad (63)$$

where $t_i(\tau_k)$ is the service time for packet i in the queue, and $t_c(\tau_k)$ is the service time for packet already in service. Also, t_s is the service time for the packet that has just arrived. we seek to find $\Pr(d(\tau_k, m) > D_i^{\max})$ which can be written as

$$\begin{aligned} \Pr(d(\tau_k, m) > D_i^{\max}) &= \sum_{m,k} \Pr(D > D_i^{\max} | m, \tau_k) \Pr(m, \tau_k) \\ &= \sum_{m,k} \Pr(D > D_i^{\max} | m, \tau_k) \Pr(m | \tau_k) \Pr(\tau_k). \end{aligned} \quad (64)$$

The probability that there are m users in an M/M/1 queue at time slot τ_k , i.e. $\Pr(m | \tau_k)$, can be written as (see [49]): $\Pr(m | \tau_k) = \left(\frac{a_i(\tau_k)}{r_i(\tau_k)} \right)^m \left(1 - \frac{a_i(\tau_k)}{r_i(\tau_k)} \right)$. Since we assumed the time slots have

equal lengths, the packets arrive at each time slot with equal probability of $\Pr(\tau_k) = \frac{1}{t}$, where t is the total number of time slots.

The sum of $m + 1$ identically independent exponential random variables with the mean $\frac{1}{r_i(\tau_k)}$ is a gamma random variable. Consequently, if the users arrive at time slot τ_k while there are m users in the system at the time of arrival, the distribution of delay is

$$f_D(\phi|m, \tau_k) = \frac{r_i(\tau_k)^{m+1}}{\Gamma(m+1)} \phi^m e^{-r_i(\tau_k)\phi}. \quad (65)$$

As a result, we can write the probability of delay exceeding a threshold $D_i^{\max}$ as

$$\Pr(D > D_i^{\max}) = \int_{D_i^{\max}}^{\infty} \sum_{m,k} f_D(\phi|m, \tau_k) \Pr(m|\tau_k) \Pr(\tau_k) d\phi \quad (66)$$

$$= \int_{D_i^{\max}}^{\infty} \frac{1}{t} \sum_{m,k} \frac{r_i(\tau_k)^{m+1}}{m!} \phi^m e^{-r_i(\tau_k)\phi} \left(\frac{a_i(\tau_k)}{r_i(\tau_k)}\right)^m \left(1 - \frac{a_i(\tau_k)}{r_i(\tau_k)}\right) d\phi \quad (67)$$

$$= \frac{1}{t} \sum_{k=1}^t \int_{D_i^{\max}}^{\infty} (r_i(\tau_k) - a_i(\tau_k)) e^{-r_i(\tau_k)\phi} \sum_{m=0}^{\infty} \frac{(\phi a_i(\tau_k))^m}{m!} d\phi \quad (68)$$

$$= \frac{1}{t} \sum_{k=1}^t \int_{D_i^{\max}}^{\infty} (r_i(\tau_k) - a_i(\tau_k)) e^{-(r_i(\tau_k) - a_i(\tau_k))\phi} d\phi \quad (69)$$

$$= \frac{1}{t} \sum_{k=1}^t e^{-(r_i(\tau_k) - a_i(\tau_k))D_i^{\max}}, \quad (70)$$

which proves the theorem.

REFERENCES

- [1] A. T. Z. Kasgari, W. Saad, and M. Debbah, "Brain-aware wireless networks: Learning and resource management," in *Proc. 51th Asilomar Conference on Signals, Systems and Computers, Pacific Grove, CA, USA*, Nov 2017.
- [2] W. Saad, M. Bennis, and M. Chen, "A vision of 6G wireless systems: Applications, trends, technologies, and open research problems," *arXiv preprint arXiv:1902.10265*, 2019.
- [3] E. Gobetti and R. Scateni, "Virtual reality: Past, present, and future," *Virtual environments in clinical psychology and neuroscience: Methods and techniques in advanced patient-therapist interaction*, Nov 1998.
- [4] M. Chen, W. Saad, and C. Yin, "Virtual reality over wireless networks: Quality-of-service model and learning-based resource management," *IEEE Transactions on Communications*, vol. 66, no. 11, pp. 5621–5635, Nov 2018.
- [5] O. Semiari, W. Saad, S. Valentin, M. Bennis, and H. V. Poor, "Context-aware small cell networks: How social metrics improve wireless resource allocation," *IEEE Transactions on Wireless Communications*, vol. 14, no. 11, pp. 5927–5940, Nov 2015.
- [6] H. Intraub, "Rapid conceptual identification of sequentially presented pictures," *Journal of Experimental Psychology: Human Perception and Performance*, vol. 7, no. 3, pp. 604–610, 1981.
- [7] K. Ur Rehman Laghari, R. Gupta, S. Arndt, J. N. Antons, R. Schleicher, S. Möller, and T. H. Falk, "Neurophysiological experimental facility for quality of experience (QoE) assessment," in *Proc. of IFIP/IEEE International Symposium on Integrated Network Management (IM)*, Ghent, Belgium, May 2013, pp. 1300–1305.
- [8] I. Wechsung and K. De Moor, "Quality of experience versus user experience," in *Quality of Experience: Advanced Concepts, Applications and Methods*, S. Möller and A. Raake, Eds. Springer International Publishing, 2014, pp. 35–54.
- [9] T. Zhao, Q. Liu, and C. W. Chen, "QoE in video transmission: A user experience-driven strategy," *IEEE Communications Surveys & Tutorials*, vol. 19, no. 1, pp. 285–302, First quarter 2017.
- [10] Y. Chen, K. Wu, and Q. Zhang, "From QoS to QoE: A tutorial on video quality assessment," *IEEE Communications Surveys & Tutorials*, vol. 17, no. 2, pp. 1126–1165, Second quarter 2015.
- [11] A. Rahmati, X. He, I. Guvenc, and H. Dai, "Dynamic mobility-aware interference avoidance for aerial base stations in cognitive radio networks," *arXiv preprint arXiv:1901.02613*, 2019.

- [12] M. Mozaffari, A. Taleb Zadeh Kargari, W. Saad, M. Bennis, and M. Debbah, "Beyond 5G with uavs: Foundations of a 3D wireless cellular network," *IEEE Transactions on Wireless Communications*, vol. 18, no. 1, pp. 357–372, Jan 2019.
- [13] Z. Zhou, H. Yu, C. Xu, Y. Zhang, S. Mumtaz, and J. Rodriguez, "Dependable content distribution in D2D-based cooperative vehicular networks: A big data-integrated coalition game approach," *IEEE Transactions on Intelligent Transportation Systems*, vol. 19, no. 3, pp. 953–964, March 2018.
- [14] A. Ferdowsi, U. Challita, W. Saad, and N. B. Mandayam, "Robust deep reinforcement learning for security and safety in autonomous vehicle systems," *Proc. of International Conference on Intelligent Transportation Systems, Maui, HI, USA*, Nov. 2018.
- [15] Z. Zhou, C. Gao, C. Xu, Y. Zhang, S. Mumtaz, and J. Rodriguez, "Social big-data-based content dissemination in internet of vehicles," *IEEE Transactions on Industrial Informatics*, vol. 14, no. 2, pp. 768–777, Feb 2018.
- [16] Z. Zhou, H. Liao, B. Gu, K. M. S. Huq, S. Mumtaz, and J. Rodriguez, "Robust mobile crowd sensing: When deep learning meets edge computing," *IEEE Network*, vol. 32, no. 4, pp. 54–60, July 2018.
- [17] M. Alam, D. Yang, K. Huq, F. Saghezchi, S. Mumtaz, and J. Rodriguez, "Towards 5G: Context aware resource allocation for energy saving," *Jour. of Signal Proc. Systems*, vol. 83, no. 2, pp. 279–291, May 2016.
- [18] Y. Lin, R. Zhang, C. Li, L. Yang, and L. Hanzo, "Graph-based joint user-centric overlapped clustering and resource allocation in ultradense networks," *IEEE Transactions on Vehicular Technology*, vol. 67, no. 5, pp. 4440–4453, May 2018.
- [19] J. Zhao, Y. Liu, K. K. Chai, M. ElKashlan, and Y. Chen, "Matching with peer effects for context-aware resource allocation in D2D communications," *IEEE Communications Letters*, vol. 21, no. 4, pp. 837–840, April 2017.
- [20] M. Zalgout, S. Abdul-Nabi, A. Khalil, M. Helard, and M. Crussiere, "Optimizing context-aware resource and network assignment in heterogeneous wireless networks," in *Proc. of IEEE Wireless Communications and Networking Conference (WCNC), San Francisco, CA, USA*, March 2017, pp. 1–6.
- [21] E. Bastug, M. Bennis, and M. Debbah, "Living on the edge: The role of proactive caching in 5G wireless networks," *IEEE Communications Magazine*, vol. 52, no. 8, pp. 82–89, Aug 2014.
- [22] P. Makris, D. N. Skoutas, and C. Skianis, "A survey on context-aware mobile and wireless networking: On networking and computing environments' integration," *IEEE Communications Surveys & Tutorials*, vol. 15, no. 1, pp. 362–386, First 2013.
- [23] Y. Nijsure, Y. Chen, C. Yuen, and Y. H. Chew, "Location-aware spectrum and power allocation in joint cognitive communication-radar networks," in *2011 6th International ICST Conference on Cognitive Radio Oriented Wireless Networks and Communications (CROWNCOM)*, June 2011, pp. 171–175.
- [24] M. Proebster, M. Kaschub, and S. Valentin, "Context-aware resource allocation to improve the quality of service of heterogeneous traffic," in *Proc. of IEEE International Conference on Communications*, Kyoto, Japan, Jul. 2011.
- [25] C. Perera, A. Zaslavsky, P. Christen, and D. Georgakopoulos, "Context aware computing for the Internet of Things: A survey," *IEEE Communications Surveys & Tutorials*, vol. 16, no. 1, pp. 414–454, First quarter 2014.
- [26] L. Huang, "System intelligence: Model, bounds and algorithms," in *Proc. of the 17th ACM International Symposium on Mobile Ad Hoc Networking and Computing*. New York, NY, USA: ACM, 2016, pp. 171–180.
- [27] H. Modares, I. Ranatunga, F. L. Lewis, and D. O. Popa, "Optimized assistive humanrobot interaction using reinforcement learning," *IEEE Transactions on Cybernetics*, vol. 46, no. 3, pp. 655–667, March 2016.
- [28] B. Qian, X. Wang, N. Cao, Y. Jiang, and I. Davidson, "Learning multiple relative attributes with humans in the loop," *IEEE Transactions on Image Processing*, vol. 23, no. 12, pp. 5573–5585, Dec 2014.
- [29] Q. Wu and R. Zhang, "Delay-constrained throughput maximization in UAV-enabled OFDM systems," in *Proc. of 23rd Asia-Pacific Conference on Communications (APCC), Perth, WA, Australia*, Dec 2017, pp. 1–6.
- [30] C. Bishop, *Pattern Recognition and Machine Learning*. New York, NY, USA: Springer Information Science and Statistics, 2007.
- [31] M. Chen, U. Challita, W. Saad, C. Yin, and M. Debbah, "Machine learning for wireless networks with artificial intelligence: A tutorial on neural networks," *arXiv preprint arXiv:1710.02913*, Oct. 2017.
- [32] S. Petkoski, A. Spiegler, T. Proix, and V. Jirsa, "Effects of multimodal distribution of delays in brain network dynamics," *BMC Neuroscience*, vol. 16, no. Suppl 1, p. P109, 2015.
- [33] Y. Yang, L. T. Park, N. B. Mandayam, I. Seskar, A. L. Glass, and N. Sinha, "Prospect pricing in cognitive radio networks," *IEEE Transactions on Cognitive Communications and Networking*, vol. 1, no. 1, pp. 56–70, Oct. 2015.
- [34] M. G. Baydogan and G. Runger, "Learning a symbolic representation for multivariate time series classification," *Data Mining and Knowledge Discovery*, vol. 29, no. 2, pp. 400–422, Mar 2015. [Online]. Available: <https://doi.org/10.1007/s10618-014-0349-y>
- [35] T. K. Moon, "The expectation-maximization algorithm," *IEEE Signal Processing Magazine*, vol. 13, no. 6, pp. 47–60, Nov 1996.
- [36] A. Jaimes and N. Sebe, "Multimodal human–computer interaction: A survey," *Computer vision and image understanding*, vol. 108, no. 1-2, pp. 116–134, Oct. 2007.
- [37] G. McLachlan and T. Krishnan, "The EM algorithm and extensions," vol. 382, 03 1998.
- [38] J. Friedman, T. Hastie, and R. Tibshirani, *The elements of statistical learning*. Springer series in statistics New York, 2001, vol. 1, no. 10.
- [39] R. Hari, S. Levnen, and T. Raij, "Timing of human cortical functions during cognition: role of meg," *Trends in Cognitive Sciences*, vol. 4, no. 12, pp. 455 – 462, 2000.

- [40] M. D. Fox and M. E. Raichle, "Spontaneous fluctuations in brain activity observed with functional magnetic resonance imaging," *Nature reviews neuroscience*, vol. 8, no. 9, p. 700, 2007.
- [41] H. Laufs, K. Krakow, P. Sterzer, E. Eger, A. Beyerle, A. Salek-Haddadi, and A. Kleinschmidt, "Electroencephalographic signatures of attentional and cognitive default modes in spontaneous brain activity fluctuations at rest," *Proceedings of the national academy of sciences*, vol. 100, no. 19, pp. 11 053–11 058, 2003.
- [42] M. J. Neely, "Stochastic network optimization with application to communication and queueing systems," *Synthesis Lectures on Communication Networks*, vol. 3, no. 1, pp. 1–211, 2010.
- [43] M. J. Neely, E. Modiano, and C. P. Li, "Fairness and optimal stochastic control for heterogeneous networks," *IEEE/ACM Transactions on Networking*, vol. 16, no. 2, pp. 396–409, April 2008.
- [44] K. Seong, M. Mohseni, and J. M. Cioffi, "Optimal resource allocation for ofdma downlink systems," in *Proc. of IEEE International Symposium on Information Theory, Seattle, WA, USA*, July 2006, pp. 1394–1398.
- [45] W. Yu and R. Lui, "Dual methods for nonconvex spectrum optimization of multicarrier systems," *IEEE Transactions on Communications*, vol. 54, no. 7, pp. 1310–1322, July 2006.
- [46] S. Boyd and A. Mutapcic, "Subgradient methods," *Lecture notes of EE364b, Stanford University, Winter Quarter*, vol. 2007, 2006.
- [47] M. Chen, W. Saad, and C. Yin, "Resource management for wireless virtual reality: Machine learning meets multi-attribute utility," in *In Proc. of IEEE Global Communications Conference (Globecom), Singapore*, Dec 2017, pp. 1–7.
- [48] J. Berger, "A robust generalized Bayes estimator and confidence region for a multivariate normal mean," *The Annals of Statistics*, pp. 716–761, 1980.
- [49] A. Papoulis and S. U. Pillai, *Probability, random variables, and stochastic processes*. Tata McGraw-Hill Education, 2002.