

Performance Optimization for Semantic Communications: An Attention-based Learning Approach

Yining Wang*, Mingzhe Chen[†], Walid Saad[‡], Tao Luo*, Shuguang Cui[§], and H. Vincent Poor[†]

*Beijing Laboratory of Advanced Information Network, Beijing University of Posts and Telecommunications, Beijing, China.

[†]Department of Electrical and Computer Engineering, Princeton University, Princeton, NJ, USA.

[‡]Wireless@VT, Bradley Department of Electrical and Computer Engineering, Virginia Tech, Blacksburg, VA, USA.

[§]Shenzhen Research Institute of Big Data and Future Network of Intelligence Institute, the Chinese University of Hong Kong, Shenzhen, China.

Emails: {wyy0206, tluo}@bupt.edu.cn, {mingzhec, poor}@princeton.edu, walids@vt.edu, shuguangcui@cuhk.edu.cn.

Abstract—In this paper, a semantic communication framework is proposed for wireless networks. In the proposed framework, a base station (BS) extracts the semantic information from textual data, and, transmits it to each user. This semantic information is modeled by a knowledge graph (KG) and hence, the semantic information consists of a set of semantic triples. After receiving the semantic information, each user recovers the original text using a graph-to-text generation model. To measure the performance of the studied semantic communication system, a metric of semantic similarity (MSS) that jointly captures the semantic accuracy and completeness of the recovered text is proposed. Due to wireless resource limitations, the BS can only transmit partial semantic information to each user so as to satisfy the transmission delay constraint. Hence, the BS must select an appropriate resource block for each user and determine partial semantic information to be transmitted. This problem is formulated as an optimization problem whose goal is to maximize the total MSS by optimizing the resource allocation policy and determining the partial semantic information to be transmitted. To solve this problem, a policy gradient-based reinforcement learning (RL) algorithm integrated with the attention network is proposed. The proposed algorithm can evaluate the importance of each triple in the semantic information using an attention network and then, build a relationship between the importance distribution of the triples in the semantic information and the total MSS. Simulation results demonstrate that the proposed semantic communication framework can reduce the size of data that the BS needs to transmit by up to 46% and yield a two-fold improvement in the total MSS compared to a standard communication network that does not consider semantic communications.

I. INTRODUCTION

The emergence of new wireless applications such as tactile Internet, interactive hologram, and intelligent humanoid robot is generating unprecedented amounts of data (at zetta-bytes scale) that will strain the capacity of modern-day wireless networks [1]. In order to support these human-centered services and applications, wireless networks must be carefully designed based on the contents, human-related requirements, human-related knowledge, and experience-based metrics [2]. These challenges can be addressed by a novel paradigm, called *semantic communication*, which allows the meaning of the data (behind digital bits) to be extracted and exploited during communication. Therefore, semantic communication has recently attracted significant interest due to its naturally advantages in terms of providing human-oriented services

and improving communication efficiency. However, deploying semantic communications in wireless networks faces several challenges including the extraction of the data semantics, semantic-oriented resource allocation, and the measurement of semantic communication performance.

Recently, several works in [3]–[6] studied a number of problems related to semantic communications. The authors in [3] proposed a deep learning approach to semantic communications that seeks to maximize the mutual information between the original and decoded signals. In [4], the authors developed a distributed semantic communication system for capacity-limited Internet of Things devices. The works in [3] and [4] considered the features extracted by deep learning models as the meaning of the data. However, the features extracted by deep learning models are unexplainable and do not have any physical meaning. Therefore, they may not be able to represent the actual meaning of the data. The authors in [5] provided an overview on the use of semantic detection and knowledge modeling technology for semantic information extraction. The work in [6] optimized the policies of transmitting the meaning of the codewords over a noisy channel so as to minimize the error between the intended meaning and the recovered message. However, none of these existing works in [3]–[6] designed a mathematical model for semantic communication-driven wireless networks. Meanwhile, they did not consider the design of a performance metric that can jointly capture the semantic communication performance and wireless communication performance.

The main contribution of this work is a novel framework that enables a wireless base station (BS) to communicate with the users using semantic communication techniques. The key contributions are summarized as follows:

- We proposed a semantic communication framework that enables a BS to transmit the meaning of the text data to its associated users. The meaning of the text data is defined as the semantic information and modeled by a knowledge graph (KG). Based on the received semantic information, the users can recover the original text using a graph-to-text generation model.
- To measure the performance of semantic communications, we propose a mathematical metric of semantic similarity (MSS) that jointly captures the semantic accuracy and completeness of the recovered text.
- Due to wireless resource limitations, the BS must de-

termine the partial semantic information (i.e., a subset of semantic triples) to be transmitted and optimize the resource allocation for each user so as to satisfy the delay constraint. The problem is formulated as an optimization problem whose goal is to maximize the total MSS by optimizing the resource allocation and determining the partial semantic information to be transmitted.

- To solve this problem, we propose an attention policy gradient (APG) algorithm that can evaluate the importance of each triple in the semantic information. Then, the proposed algorithm can analyze the relationship between the importance distribution of the triples in the semantic information and the total MSS, thus finding the effective policies for resource allocation and semantic information transmission.

Simulation results show that, compared to a standard communication network that does not consider semantic communication, the proposed APG algorithm can reduce 46% size of data that the BS needs to transmit and yield a two-fold improvement in the total MSS. To our knowledge, *this is the first work that introduces a mathematical model for semantic communication enabled wireless networks and optimizes resource allocation and semantic information transmission to improve the performance of semantic-driven wireless networks.*

II. SYSTEM MODEL AND PROBLEM FORMULATION

Consider a cellular network in which a BS transmits text data to a set $\mathcal{U}$ of U users using semantic communication techniques. Semantic communication techniques enable the BS to transmit the meaning of the text to each user so as to reduce the size of the data transmitted over wireless links. Hereinafter, the meaning of the text transmitted over wireless links is called *semantic information*. To achieve semantic communication, the BS must extract the semantic information from the original text and send it to the corresponding user¹. Then, each user recovers the text based on the received semantic information. In particular, the procedure of the considered semantic communication consists of three phases (shown in Fig. 1): a) semantic information extraction, b) semantic information transmission, and c) original text data recovery. Next, we first introduce the process of the semantic communications. Then, we introduce a semantic similarity model to measure the quality of semantic communications.

A. Semantic Information Extraction

In the studied model, we use $w_{i,n}$ to represent a word, a symbol, or a punctuation in the text data. Hereinafter, $w_{i,n}$ is called a token. Hence, the text data that the BS needs to transmit to user i consists of a sequence of tokens, as follows:

$$L_i = \{w_{i,1}, w_{i,2}, \dots, w_{i,n}, \dots, w_{i,N_i}\}, \forall w_{i,n} \in \mathcal{V}, \quad (1)$$

where $\mathcal{V}$ is the vocabulary and N_i is the number of tokens in L_i . For example, we assume that the text data required to

¹In this work, we only consider the text data transmission. One can easily extend the proposed model to other types of data such as audio data and image data.

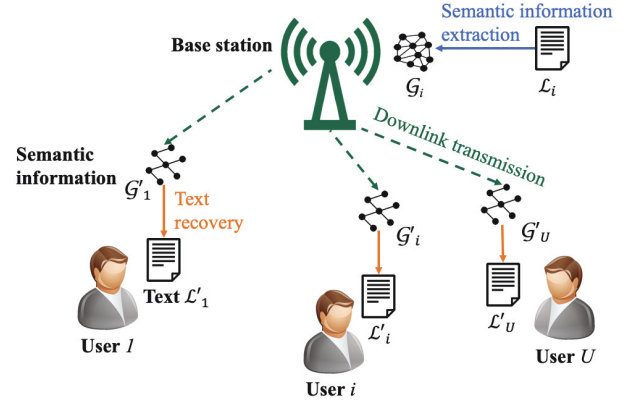


Fig. 1: The semantic communication technique-enabled wireless network.

transmit to user i is “How are you?”. Hence, we have $L_i = \{\text{[how]}, \text{[are]}, \text{[you]}, \text{[?]} \}$, where $w_{i,1} = \text{[how]}$, $w_{i,2} = \text{[are]}$, $w_{i,3} = \text{[you]}$, and $w_{i,4} = \text{[?]}$.

In our model, the semantic information extracted from a text data is modeled by a KG [7]. Hence, the semantic information consists of a set of nodes and a set of edges, as shown in Fig. 2. In particular, each node in the semantic information is an *entity* that refers to an object or a concept in the real word. Hereinafter, we define entity j in text L_i as $e_{i,j}$ that consists of a subsequence of tokens in text L_i . For example, in Fig. 2, “*stochastic lexicon model*” is an entity that consists of three tokens in the original text. An information extraction system such as the scientific information extractor in [8] can be used to recognize the set $\mathcal{E}_i$ of E_i entities in text L_i .

Edges are the *relations* between each pair of entities. Given a pair of recognized entities $(e_{i,j}, e_{i,k}), j \neq k$, the BS must find the relation $r_{i,jk} \in \mathcal{R}_i$ between them, where $\mathcal{R}_i$ is the set of R_i relations involved in text L_i . For example, in Fig. 2, the relation between entity “*stochastic lexicon model*” and “*speech recognizer*” can be formulated as “*used for*”. Note that the relations (i.e., the edges of the semantic information) are directional and hence, we have $r_{i,jk} \neq r_{i,kj}$. We assume that the set $\mathcal{R}$ that includes all relations in the texts of all users is predefined and each relation is a two-token sequence such as “*part of*” and “*used for*”, as done in [8]. Hence, given $e_{i,j}$ and $e_{i,k}$ in the original text L_i , the relation $r_{i,jk}$ between $e_{i,j}$ and $e_{i,k}$ can be obtained by classification algorithms such as convolutional neural networks [9].

Based on the recognized entities and the extracted relations, the semantic information of text L_i can be expressed as

$$\mathcal{G}_i = \{\varepsilon_i^1, \dots, \varepsilon_i^g, \dots, \varepsilon_i^{G_i}\}, \quad (2)$$

where $\varepsilon_i^g = (e_{i,j}^g, r_{i,jk}^g, e_{i,k}^g), \forall e_{i,j}^g, e_{i,k}^g \in \mathcal{E}_i, j \neq k, \forall r_{i,jk}^g \in \mathcal{R}_i$ is a semantic triple and G_i is the number of semantic triples in $\mathcal{G}_i$. Since either each entity (e.g., $e_{i,j}^g$ or $e_{i,k}^g$) or each relation $r_{i,jk}^g$ consists of a sequence of tokens (e.g., $e_{i,j} = \{\text{[stochastic]}, \text{[lexicon]}, \text{[model]}\}$, $e_{i,k} = \{\text{[speech]}, \text{[recognizer]}\}$, and $r_{i,jk} = \{\text{[used]}, \text{[for]}\}$), triple ε_i^g can be expressed as $\varepsilon_i^g = \{v_{i,1}^g, \dots, v_{i,b}^g, \dots, v_{i,B^g}^g\}$, where

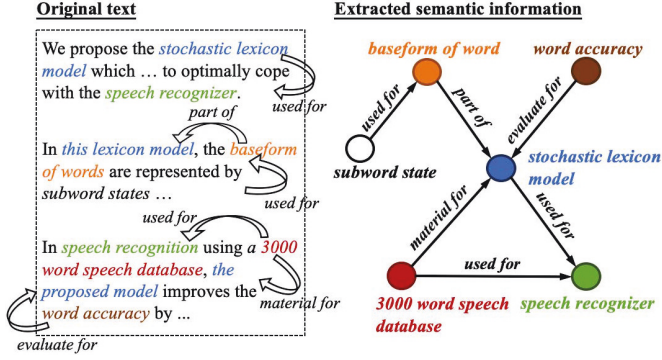


Fig. 2: An example of an original text and the extracted semantic information.

$v_{i,b}^g \in \mathcal{V}$ is token b in semantic triple ε_i^g and $B_i^g = S_{i,j}^g + S_{i,k}^g + 2$ with $S_{i,j}^g$ being the number of tokens in entity $e_{i,j}^g$. The number of tokens in semantic information $\mathcal{G}_i$ is

$$Z(\mathcal{G}_i) = \sum_{g=1}^{G_i} (S_{i,j}^g + S_{i,k}^g + 2). \quad (3)$$

From Fig. 2, we see that the data size of the extracted semantic information is much smaller than the data size of the original text (i.e., $Z(\mathcal{G}_i) \ll N_i$). This is because the two-token relations between the entity pairs in $\mathcal{G}_i$ can reduce the redundant context in original text L_i .

B. Transmission Model

We assume that an orthogonal frequency division multiple access (OFDMA) technique is adopted. The BS has a set $\mathcal{Q}$ of $Q \leq U$ downlink orthogonal resource blocks (RBs) that can be allocated to serve the users. The RB allocation vector of user i is $\alpha_i = [\alpha_{i,1}, \dots, \alpha_{i,q}, \dots, \alpha_{i,Q}]$, where $\alpha_{i,q} \in \{0, 1\}$. Here, $\alpha_{i,q} = 1$ implies that RB q is allocated to user i ; otherwise, we have $\alpha_{i,q} = 0$. In our model, we assume that each user can only occupy one RB and each RB can only be allocated to one user [10]. Then, we have

$$\sum_{q=1}^Q \alpha_{i,q} \leq 1, \forall i \in \mathcal{U}; \quad \sum_{i=1}^U \alpha_{i,q} \leq 1, \forall q \in \mathcal{Q}. \quad (4)$$

The downlink channel capacity of the BS transmitting semantic information $\mathcal{G}_i$ to user i is given as

$$c_i(\alpha_i) = \sum_{q=1}^Q \alpha_{i,q} W \log_2 \left(1 + \frac{Pr_i}{I_q + WN_0} \right) \quad (5)$$

where W is the bandwidth of each RB, P is the transmit power of the BS, I_q represents the interference caused by BSs that are located in other service areas and use RB q , N_0 is the noise power spectral density, and $r_i = \gamma_i d_i^{-2}$ is the channel gain between the BS and user i with γ_i being the Rayleigh fading parameters and d_i being the distance between the BS and user i . We assume that the transmission delay between the BS and each user i is limited to T . Hence, the BS must select partial semantic information (i.e., a subset of triples) to transmit so as to satisfy the transmission delay constraint.

C. Text Recovery

The partial semantic information that the BS needs to transmit to user i is given as

$$\mathcal{G}'_i = \{\varepsilon_i^1, \dots, \varepsilon_i^h, \dots, \varepsilon_i^{H_i}\} \subset \mathcal{G}_i, \quad (6)$$

where $\varepsilon_i^h = (e_{i,j}^h, r_{i,j,k}^h, e_{i,k}^h)$ and H_i is the number of selected semantic triples in $\mathcal{G}'_i$. Given the transmission delay threshold T , the selected semantic information $\mathcal{G}'_i$ should satisfy the delay constraint as follows:

$$\frac{Z(\mathcal{G}'_i)R}{c_i(\alpha_i)} \leq T, \quad (7)$$

where $Z(\mathcal{G}'_i) = \sum_{h=1}^{H_i} (S_{i,j}^h + S_{i,k}^h + 2)$ and R is the number of bits used to represent each token. After each user i receives the semantic information $\mathcal{G}'_i$, a graph-to-text generation model, such as the graph transformer in [11], can be used to recover the coherent multi-sentence text from $\mathcal{G}'_i$. We assume that the graph-to-text generation model [11] is well-trained and shared among all users. The text recovered by user i based on $\mathcal{G}'_i$ is

$$L'_i(\alpha_i, \mathcal{G}'_i) = \{w'_{i,0}, w'_{i,1}, \dots, w'_{i,m}, \dots, w'_{i,M_i}\}, \quad (8)$$

where M_i is the number of tokens in the recovered text L'_i .

D. Semantic Similarity Model

To measure the quality of the semantic communication, we proposed a metric of semantic similarity (MSS). Different from the existing metric of BLEU [12] that is only focus on the semantic accuracy of the recovered text, the proposed MSS jointly capture the semantic accuracy and completeness of the recovered text compared to the original text. The semantic accuracy of recovered text $L'_i(\alpha_i, \mathcal{G}'_i)$ is defined as [12]

$$A_i(\alpha_i, \mathcal{G}'_i) = \frac{\sum_{m=1}^{M_i} \min(\sigma(L'_i(\alpha_i, \mathcal{G}'_i), w'_{i,m}), \sigma(L_i, w'_{i,m}))}{\sum_{m=1}^{M_i} \sigma(L'_i(\alpha_i, \mathcal{G}'_i), w'_{i,m})}, \quad (9)$$

where $\sigma(L'_i(\alpha_i, \mathcal{G}'_i), w'_{i,m})$ is the number of occurrences of $w'_{i,m}$ in recovered text $L'_i(\alpha_i, \mathcal{G}'_i)$, $\sigma(L_i, w'_{i,m})$ is the number of occurrences of $w'_{i,m}$ in original text L_i . The semantic completeness of the recovered text is defined as [13]

$$R_i(\alpha_i, \mathcal{G}'_i) = \frac{\sum_{m=1}^{M_i} \min(\sigma(L'_i(\alpha_i, \mathcal{G}'_i), w'_{i,m}), \sigma(L_i, w'_{i,m}))}{\sum_{m=1}^{M_i} \sigma(L_i, w'_{i,m})}. \quad (10)$$

Next, we use an example to explain the differences between the semantic accuracy and the semantic completeness more clearly. For example, $L_i = \{[\text{how}], [\text{are}], [\text{you}], [?]\}$ and $L'_i(\alpha_i, \mathcal{G}'_i) = \{[\text{how}], [\text{do}], [\text{you}], [\text{do}], [?]\}$. We have $\sigma(L_i, [\text{how}]) = 1$, $\sigma(L'_i(\alpha_i, \mathcal{G}'_i), [\text{how}]) = 1$, $\sigma(L_i, [\text{do}]) = 0$, and $\sigma(L'_i(\alpha_i, \mathcal{G}'_i), [\text{do}]) = 2$. From (9) and (10), we can obtain $A_i(\alpha_i, \mathcal{G}'_i) = \frac{1+0+1+0+1}{1+2+1+2+1} = \frac{3}{7}$ and $R_i(\alpha_i, \mathcal{G}'_i) = \frac{1+0+1+0+1}{1+0+1+0+1} = 1$, respectively. Based on (9) and (10), the MSS of recovered text $L'_i(\alpha_i, \mathcal{G}'_i)$ can be given as

$$E_i(\alpha_i, \mathcal{G}'_i) = \theta_i \frac{A_i(\alpha_i, \mathcal{G}'_i)R_i(\alpha_i, \mathcal{G}'_i)}{\varphi A_i(\alpha_i, \mathcal{G}'_i) + (1-\varphi)R_i(\alpha_i, \mathcal{G}'_i)}, \quad (11)$$

where $\varphi \in (0, 1)$ is a parameter used to adjust the contributions of semantic accuracy and completeness to the MSS. In particular, increasing the value of φ will increase the effect of semantic accuracy on the MSS. θ_i is an additional penalty for short text and can be represented by [12]

$$\theta_i = \begin{cases} 1, & M_i \geq N_i, \\ e^{1 - \frac{N_i}{M_i}}, & M_i < N_i. \end{cases} \quad (12)$$

Using θ_i , the recovered text with more tokens can result in a higher MSS. From (11), we see that the proposed MSS used to evaluate the recovered text can control the tradeoff between the semantic accuracy and the semantic completeness. In particular, for each token $w'_{i,m}$ in the recovered text $L'_i(\alpha_i, \mathcal{G}'_i)$, if it appears more times in $L'_i(\alpha_i, \mathcal{G}'_i)$ than L_i , then $A_i(\alpha_i, \mathcal{G}'_i)$ decreases; otherwise, $R_i(\alpha_i, \mathcal{G}'_i)$ decreases.

E. Problem Formulation

Given the defined system model, our goal is to maximize the total MSS of all the texts recovered by the users while satisfying the transmission delay requirement. This maximization problem includes optimizing the RB allocation and determining the partial semantic information transmitted to each user. The MSS maximization problem is formulated as follows:

$$\max_{\alpha_i, \mathcal{G}'_i} \sum_{i=1}^U E_i(\alpha_i, \mathcal{G}'_i), \quad (13)$$

$$\text{s.t. } \alpha_{i,q} \in \{0, 1\}, \forall i \in \mathcal{U}, \forall q \in \mathcal{Q}, \quad (13a)$$

$$\sum_{q=1}^Q \alpha_{i,q} \leq 1, \forall i \in \mathcal{U}, \quad (13b)$$

$$\sum_{i=1}^U \alpha_{i,q} \leq 1, \forall q \in \mathcal{Q}, \quad (13c)$$

$$\frac{Z_i(\mathcal{G}'_i)R}{c_i(\alpha_i)} \leq T, \forall i \in \mathcal{U}, \quad (13d)$$

where (13a), (13b), and (13c) guarantee that each user can only occupy one RB and each RB can only be allocated to one user for semantic information transmission. (13d) is the delay requirement of semantic information transmission. From (13), we see that the MSS jointly depends on the selected subset $\mathcal{G}'_i$ of the semantic information and the RB allocation α_i . However, the objective function of problem (13) is non-convex and depends on the text generation model used to recover the text. Hence, (13) cannot be solved by the traditional optimization algorithms. To solve (13), we use a reinforcement learning (RL) algorithm [14] with an attention network that can evaluate the importance of each semantic triple so as to build the relationship between the importance of the triples in the semantic information and the total MSS of the recovered texts.

III. ATTENTION RL FOR SEMANTIC INFORMATION SELECTION AND RESOURCE ALLOCATION

Next, we introduce a policy gradient-based RL algorithm [15] integrated with an attention network [16], called attention

policy gradient (APG), that can effectively solve problem (13). First, for each semantic triple ε_i^g , we use an attention network to calculate its corresponding importance value. Here, the importance of a semantic triple is defined as the correlation between the triple and the original text. Hereinafter, we use an importance vector $\mathbf{f}_i(\mathcal{G}_i)$ to represent the *importance distribution* of the triples in each semantic information $\mathcal{G}_i$. Based on the importance evaluation, the proposed APG algorithm enables the BS to analyze the relationship between the importance distribution $\mathbf{f}_i(\mathcal{G}_i)$ and the total MSS so as to optimize the semantic information selection and RB allocation for the total MSS improvement. Next, we first introduce the use of an attention network to calculate the importance of the triples in each semantic information. Then, we show the components of the APG algorithm. Finally, we explain the entire procedure of using our APG algorithm to determine the partial semantic information to be transmitted and optimize RB allocation for each user.

A. Attention Network for Importance Evaluation

The input of an attention network is a triple ε_i^g and the original text L_i . To obtain the importance of each triple ε_i^g , the BS first needs to vectorize the tokens in ε_i^g and L_i as done in [16]. Then, we define the vector used to represent token $v_{i,b}^g$ as $\mathbf{x}_{i,b}^g \in \mathbb{R}^{D_x}$ and the vector used to represent token $w_{i,n}$ as $\mathbf{x}_{i,n} \in \mathbb{R}^{D_x}$, where D_x is the dimension of each token vector. Given the token vectors, the correlation between triple ε_i^g and token $w_{i,n}$ in L_i can be given as

$$\beta_i^g(w_{i,n}) = \sum_{b=1}^{B_i^g} \frac{(\mathbf{W}_{\text{tri}} \mathbf{x}_{i,b}^g)^\top (\mathbf{W}_{\text{tok}} \mathbf{x}_{i,n})}{B_i^g}, \quad (14)$$

where B_i^g is the number of tokens in triple ε_i^g , $\mathbf{W}_{\text{tri}} \in \mathbb{R}^{D_a \times D_x}$ and $\mathbf{W}_{\text{tok}} \in \mathbb{R}^{D_a \times D_x}$ are both parameter matrices of the attention network with $D_a \times D_x$ being the size of the parameter matrices. Using the trained parameters, an attention network can calculate $(\mathbf{W}_{\text{tri}} \mathbf{x}_{i,b}^g)^\top (\mathbf{W}_{\text{tok}} \mathbf{x}_{i,n})$ which indicates the correlation between token $v_{i,b}^g$ in triple ε_i^g and token $w_{i,n}$ in the original text L_i . The importance of ε_i^g that is defined as the correlation between triple ε_i^g and text L_i is given as

$$\omega_i^g = \sum_{n=1}^{N_i} \beta_i^g(w_{i,n}). \quad (15)$$

The importance distribution of semantic information $\mathcal{G}_i$ is

$$\mathbf{f}_i(\mathcal{G}_i) = \frac{e^{\omega_i}}{\sum_{g=1}^G e^{\omega_g}}, \quad (16)$$

where $\boldsymbol{\omega}_i = [\omega_i^1, \dots, \omega_i^g, \dots, \omega_i^{G_i}]$.

B. Components of APG Algorithm

An APG algorithm consists of six components: a) agent, b) environment, c) actions, d) states, e) policy, and f) reward, which are specified as follows

- *Agent*: Our agent is the BS that must determine the RB allocation and the semantic information selection.

Algorithm 1 APG algorithm for RB allocation and semantic information selection.

- 1: **Input:** Original text L_i for each user, transmission delay threshold T .
- 2: **Initialize:** Parameters θ generated randomly, interference I_q of each RB, task learning rate δ , and number of iterations E .
- 3: Calculate the importance distribution $f(\mathcal{G}_i)$ based on (16).
- 4: **for** $i = 1 \rightarrow E$ **do**
- 5: Collect D trajectories $\mathcal{D} = \{\alpha_1, \dots, \alpha_D\}$ using π_θ .
- 6: Update the parameters of the policy based on (19).
- 7: **end for**

- **Actions:** Each action of the agent is a vector $\alpha = [\alpha_1, \dots, \alpha_U]$. Here, once the BS determines the RB allocation, the semantic information $\mathcal{G}'_i$ that consists the most important triples in $\mathcal{G}_i$ while satisfying $\frac{Z_i(\mathcal{G}'_i)R}{c_i(\alpha_i)} \leq T$ can be determined. The action space $\mathcal{A}$ is the set of all optional actions that satisfy the constraints in (13).
- **States:** The state is defined as $s = [f(\mathcal{G}_1), \dots, f(\mathcal{G}_U)]$.
- **Policy:** The policy is the probability of the agent choosing each action given state s . The APG algorithm uses a deep neural network (DNN) parameterized by θ to build the relationship between the importance distributions and the total MSS. Hence, the trained DNN can map the input state to the output policy that maximizes the total MSS. Then, the policy can be expressed as $\pi_\theta(s, \alpha) = P(\alpha|s)$.
- **Reward:** The reward of choosing action α based on state s is $R(\alpha|s) = \sum_{i=1}^U E_i(\alpha_i, \mathcal{G}'_i)$ which is equivalent to the objective function of problem (13). Here, each $\mathcal{G}'_i$ is determined based on the selected action α .

C. APG for Total MSS Maximization

Next, we introduce the entire procedure of training the proposed APG algorithm for solving problem (13). To solve problem (13), the agent first samples D actions according to an initial policy π_θ . The set of collected actions is $\mathcal{D} = \{\alpha_1, \dots, \alpha_d, \dots, \alpha_D\}$. To evaluate the policy π_θ for the total MSS maximization, we define the expected reward of the actions in $\mathcal{D}$ as

$$\bar{J}(\theta) = \sum_{d=1}^D R(\alpha_d|s) \pi_\theta(s, \alpha_d). \quad (17)$$

The goal of optimizing the policy π_θ is to maximize the total MSS of the texts recovered by all users, that is

$$\max_{\theta} \bar{J}(\theta). \quad (18)$$

Then, the policy π_θ can be updated using the standard gradient ascent method

$$\theta \leftarrow \theta + \delta \nabla_{\theta} \bar{J}(\theta), \quad (19)$$

where α is the learning rate and $\nabla_{\theta} \bar{J}(\theta)$ is the gradient of parameter θ .

By iteratively running the policy updating step, the parameter θ of the policy can find the relation between the importance distributions of all semantic information and the total MSS. Hence, the policy for RB allocation and semantic information selection that can achieve maximum total MSS of all recovered texts can be obtained [17]. The specific training process of the proposed APG algorithm is summarized in **Algorithm 1**.

TABLE I: System Parameters

Parameters	Value	Parameters	Value
Q	10	W	20 MHz
P	1 W	N_0	-174 dBm/Hz
T	20 ms	R	80 bit
φ	0.5	D	100
D_a	64	D_x	500

Original text	We propose the stochastic lexicon model which represents the pronunciation variations to optimally cope with the continuous speech recognizer. In this lexicon model, the baseform of words are represented by subword states and the probability distribution of subwords as a hidden Markov model. Also, the proposed approach can be applied to a system employing non-linguistic recognition units and the lexicon is automatically trained from training utterances. In speaker independent speech recognition tests using a 3000 word continuous speech database, the proposed system improves the word accuracy by about 27.8% and the sentence accuracy by about 22.4%.	Correlation $\times 10^{-3}$ 3 2.5 2 1.5 1
	(stochastic lexicon model, used for, speech recognizer), (baseform of word, part of, stochastic lexicon model), (probability distribution of subwords, conjunction with, baseform of words), (3000 word speech database, material for, stochastic lexicon model), (word accuracy; evaluate for; stochastic lexicon model), (sentence accuracy; conjunction with; word accuracy).	
	in this paper we present a new STOCHASTIC LEXICON MODEL for SPEECH RECOGNIZER. the proposed STOCHASTIC LEXICON MODEL consists of a BASEFORM OF WORDS and a PROBABILITY DISTRIBUTION OF SUBWORDS. a STOCHASTIC LEXICON MODEL is trained on a 3000 WORD CONTINUOUS SPEECH DATABASE. the experimental results show that the proposed STOCHASTIC LEXICON MODEL achieves higher WORD ACCURACY and SENTENCE ACCURACY than the conventional STOCHASTIC LEXICON MODEL.	

Fig. 3: An example of the original text, the transmitted semantic information, and the recovered text.

IV. SIMULATION RESULTS AND ANALYSIS

In our simulations, a circular network is considered with one BS and $U = 10$ uniformly distributed users. Other parameters are listed in Table I. We use the semantic information extraction model in [8] and the text recovery model in [11]. The text dataset used to train the proposed APG algorithm is the abstract generation dataset (AGENDA) [18] that consists of 40 thousand paper titles and abstracts from the proceedings of 12 top artificial intelligence conferences. For comparison purposes, we consider two baselines: a) the traditional policy gradient algorithm that learns a policy for RB allocation and selects $\mathcal{G}'_i$ randomly, and b) the traditional symbols-based wireless communication scheme that directly transmits the original text data.

Fig. 3 shows an example of using proposed semantic communication framework for text data transmission. In particular, Fig. 3 shows the original text, the transmitted semantic information, and the recovered text. In Fig. 3, we use different colors to represent the correlation between the semantic triple “(stochastic lexicon model, used for, speech recognize)” and different tokens in the original text. In particular, as the correlations between the semantic triple “(stochastic lexicon model, used for, speech recognize)” and the tokens in the original text increase, the color used to mark the tokens changes from white to blue. In Fig. 3, the importance of the triple “(stochastic lexicon model, used for, speech recognize)” is shown as the sum of correlations. According to the order of importance of the triples, the selected partial semantic information to be transmitted is listed in Fig. 3. Fig. 3 also shows that the text recovered by the user covers the main

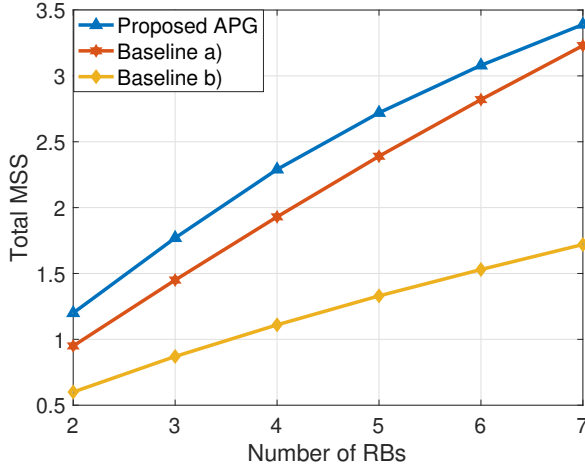


Fig. 4: The total MSS as the number of RBs varies.

meaning of the original text. This is because the BS selects the most important semantic information based on the importance evaluation and transmits it to the user. From Fig. 3, we also see that, using our proposed semantic communication framework, the BS needs to only transmit 58 tokens of the semantic information instead of 108 words of the original text data to the user. Therefore, the proposed framework can reduce 46% data transmitted over wireless links.

Fig. 4 shows how the total MSS changes as the number of RBs varies. In Fig. 4, we can see that, as the number of RBs increases, the total MSS increases. This is because as the number of RBs increases, the number of users that receive the semantic information from the BS increases. Fig. 4 also shows that the proposed APG algorithm can achieve up to 13.2% and 101.8% gains in terms of the total MSS compared to baselines a) and b). This is because the proposed APG algorithm can optimize the RB allocation and determine the partial semantic information to be transmitted.

V. CONCLUSION

In this paper, we have developed a novel semantic communication framework for wireless networks. We have modeled the meaning of the text data by a KG. To measure the performance of the semantic communications, we have introduced the MSS that captures the semantic similarity between the original text and the recovered text. We have jointly considered the wireless resource limitations and the performance of the semantic communications and formulated an optimization problem whose goal is to maximize the total MSS by optimizing the RB allocation and semantic information transmission. To solve this problem, we have developed an APG algorithm that can obtain the importance distribution of the triples in the semantic information and then build the relationship between the importance distribution and the total MSS. Hence, the proposed APG algorithm enables the BS to find the policies for RB allocation and semantic information selection for maximizing the total MSS. Simulation results have demonstrated that, compared with a standard communication network that does

not consider semantic communication, the proposed semantic communication framework can significantly reduce the size of data required to transmit and increase the total MSS.

REFERENCES

- [1] W. Saad, M. Bennis, and M. Chen, "A vision of 6G wireless systems: Applications, trends, technologies, and open research problems," *IEEE Network*, vol. 34, no. 3, pp. 134–142, May 2020.
- [2] M. Kountouris and N. Pappas, "Semantics-empowered communication for networked intelligent systems," *IEEE Communications Magazine*, vol. 59, no. 6, pp. 96–102, June 2021.
- [3] H. Xie, Z. Qin, G. Li, and B. Juang, "Deep learning enabled semantic communication systems," *IEEE Transactions on Signal Processing*, vol. 69, pp. 2663–2675, April 2021.
- [4] H. Xie and Z. Qin, "A lite distributed semantic communication system for internet of things," *IEEE Journal on Selected Areas in Communications*, vol. 39, no. 1, pp. 142–153, Jan. 2021.
- [5] G. Shi, Y. Xiao, Y. Li, and X. Xie, "From semantic communication to semantic-aware networking: model, architecture, and open problems," 2020, Available: <https://arxiv.org/abs/2012.15405>.
- [6] B. Güler, A. Yener, and A. Swami, "The semantic communication game," *IEEE Transactions on Cognitive Communications and Networking*, vol. 4, no. 4, pp. 787–802, Dec. 2018.
- [7] S. Ji, S. Pan, E. Cambria, P. Martinen, and P. S. Yu, "A survey on knowledge graphs: Representation, acquisition, and applications," *IEEE Transactions on Neural Networks and Learning Systems*, to appear 2021.
- [8] Y. Luan, L. He, M. Ostendorf, and H. Hajishirzi, "Multi-task identification of entities, relations, and coreference for scientific knowledge graph construction," in *Proc. Conference on Empirical Methods in Natural Language Processing (EMNLP)*, Brussels, Belgium, Oct. 2018, pp. 3219–3232.
- [9] Y. Lin, S. Shen, Z. Liu, H. Luan, and M. Sun, "Neural relation extraction with selective attention over instances," in *Proc. Annual Meeting of the Association for Computational Linguistics (ACL)*, Berlin, Germany, Aug. 2016, pp. 2124–2133.
- [10] M. Chen, Z. Yang, W. Saad, C. Yin, H. V. Poor, and S. Cui, "A joint learning and communications framework for federated learning over wireless networks," *IEEE Transactions on Wireless Communications*, vol. 20, no. 1, pp. 269–283, Jan. 2021.
- [11] R. Koncel-Kedziorski, D. Bekal, Y. Luan, M. Lapata, and H. Hajishirzi, "Text generation from knowledge graphs with graph transformers," in *Proc. Conference of the North American Chapter of the Association for Computational Linguistics*, Minneapolis, Minnesota, June 2019, vol. 1, pp. 2284–2293.
- [12] K. Papineni, S. Roukos, T. Ward, and W. Zhu, "BLEU: a method for automatic evaluation of machine translation," in *Proc. Annual Meeting of the Association for Computational Linguistics*, Philadelphia, PA, USA, July 2002, pp. 311–318.
- [13] S. Banerjee and A. Lavie, "METEOR: An automatic metric for MT evaluation with improved correlation with human judgments," in *Proc. Annual Meeting of the Association for Computational Linguistics Workshop*, Ann Arbor, MI, USA, June 2005, pp. 65–72.
- [14] M. Chen, D. Gündüz, K. Huang, W. Saad, M. Bennis, A. V. Feljan, and H. V. Poor, "Distributed learning in wireless networks: Recent progress and future challenges," *IEEE Journal on Selected Areas in Communications*, to appear 2021.
- [15] S. Wang, M. Chen, Z. Yang, C. Yin, W. Saad, S. Cui, and H. V. Poor, "Reinforcement learning for minimizing age of information in real-time internet of things systems with realistic physical dynamics," 2021, Available: <https://arxiv.org/abs/2104.01527>.
- [16] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, U. Kaiser, and I. Polosukhin, "Attention is all you need," in *Proc. International Conference on Neural Information Processing Systems*, Long Beach, CA, USA, Dec. 2017, vol. 11, p. 6000–6010.
- [17] Y. Hu, M. Chen, W. Saad, H. V. Poor, and S. Cui, "Distributed multi-agent meta learning for trajectory design in wireless drone networks," *IEEE Journal on Selected Areas in Communications*, to appear 2021.
- [18] W. Ammar et al., "Construction of the literature graph in semantic scholar," 2018, Available: <https://arxiv.org/abs/1805.02262>.