

FedFA: *Firefly Algorithm for Communication Cost Optimization in Federated Learning*

Article History:

Received
12 December 2024
Revised
21 January 2025
Accepted
17 February 2025

**NIFTY FATH, PEBY WAHYU PURNAWAN, DIDI
KRISTANTO, RIDHA MUFIDA**

Program Studi Teknik Elektro, Universitas Budi Luhur, Indonesia
Email: nifty.fath@budiluhur.ac.id

ABSTRAK

Federated Learning (FL) muncul sebagai model komunikasi yang menjanjikan untuk mengatasi isu keamanan dan privasi data. Setiap perangkat klien terlibat dalam model pembelajaran mesin yang kolaboratif sehingga tidak perlu mengirimkan keseluruhan data klien ke server. Namun, salah satu hambatan penerapan FL pada komunikasi jaringan nirkabel adalah keterbatasan bandwidth serta kondisi jaringan yang tidak stabil. Oleh karena itu, penelitian ini mengusulkan pendekatan baru FedFA yang mengintegrasikan algoritma Firefly dengan tujuan mengoptimasi inisialisasi bobot sehingga dapat meminimalisir communication cost. Prinsip dasar dari FedFA adalah melibatkan parameter-parameter di dalam algoritma Firefly untuk memilih bobot terbaik dari setiap klien untuk dapat dilatih di server. Berdasarkan hasil pengujian, algoritma yang diusulkan menghasilkan peningkatan akurasi sebesar 12,84% dibandingkan dengan FedAvg. Model FedFA juga lebih tahan terhadap kondisi jaringan yang tidak stabil, terlihat dari penurunan akurasi yang tidak terlalu signifikan jika dibandingkan dengan algoritma FedAvg.

Kata kunci: *algoritma firefly, Federated Learning, machine learning, communication cost, optimasi*

ABSTRACT

Federated Learning is a promising communication model to address data security and privacy issues. Each client device engages in a collaborative machine learning model, eliminating the need to send all client data to the server. However, the main obstacles to applying FL to wireless network communication are limited bandwidth and unstable network conditions. Therefore, this research proposes a new FedFA approach integrating the Firefly algorithm to optimize weight initialization and minimize communication costs. The basic principle of FedFA involves parameters in the Firefly algorithm to select the best weight of each client to be trained on the server. Based on the test results, the proposed algorithm produces an accuracy improvement of 12.84% compared to FedAvg. The FedFA model is also more resilient to unstable communication, as seen from the less significant decrease in accuracy compared to the FedAvg algorithm.

Keywords: *firefly algorithm, Federated Learning, machine learning, communication cost, optimization*

This is an open-access article under the [CC BY-SA](https://creativecommons.org/licenses/by-sa/4.0/) license



1. PENDAHULUAN

Perkembangan teknologi telekomunikasi ditandai dengan jumlah perangkat nirkabel yang terus meningkat. Berdasarkan data yang dihimpun oleh Cisco pada tahun 2023, terdapat 29,3 milyar piranti yang terkoneksi ke jaringan internet **(Cisco, 2020)**. Jumlah tersebut akan terus bertambah secara eksponensial dari tahun ke tahun. Keberagaman jenis dan peningkatan jumlah pengguna piranti komunikasi cerdas harus mampu memenuhi beberapa standar kualifikasi telekomunikasi yakni latensi yang rendah, kecepatan transfer data yang tinggi, keamanan data yang terjamin, hemat energi, serta rendah biaya **(Lu, dkk, 2021)**. Untuk memenuhi standar tersebut, beberapa aplikasi kecerdasan buatan sudah mulai banyak digunakan, salah satunya adalah *federated learning*.

Federated learning merupakan suatu model *machine learning* yang dikenalkan oleh Google pada tahun 2016 **(Konečný, dkk, 2016)**. Pada model *federated learning*, suatu server akan membagikan sebuah model global kepada klien yang berpartisipasi. Setiap klien kemudian melakukan pembelajaran lokal menggunakan data lokal yang dimiliki oleh klien. Setelah proses pembelajaran lokal selesai dilakukan, setiap klien mengunggah pembaruan model yang baru kepada server melalui jaringan nirkabel. Server akan melakukan agregasi terhadap seluruh model lokal yang diterima dan memperbarui model global. Model global yang terbaru akan dikirim kembali ke seluruh klien untuk dilakukan pembelajaran lokal kembali pada setiap klien. Proses pembaruan model terus dilakukan hingga telah mencapai konvergensi model yang optimal

Konsep pembelajaran dalam *federated learning* memungkinkan banyak piranti untuk melatih sebuah model tanpa harus membagikan data lokal ke server. Metode ini berbeda dengan metode konvensional sebelumnya yakni metode desentralisasi yang mengharuskan seluruh dataset lokal diunggah ke suatu server. Perkembangan komputasi tepi (*edge computing*) yang pesat memungkinkan model *federated learning* untuk dapat diterapkan di berbagai komunikasi nirkabel seperti jaringan 5G dan jaringan IoT **(Nguye, dkk, 2021)**.

Namun, terdapat beberapa hal yang perlu dipertimbangkan dalam menerapkan model *federated learning* ke dalam komunikasi nirkabel yakni sebagai berikut **(Li, dkk, 2020) (Niknam, dkk, 2020) (Qin, dkk, 2021) (Zhang, dkk, 2022)**.

1. *Communication cost* untuk mengelola dan mengagregasi seluruh data yang diperoleh dari setiap pembelajaran lokal cukup tinggi
2. Komunikasi nirkabel identik dengan jaringan yang tidak selalu stabil sehingga dapat mengganggu proses *transfer learning* yang dapat mengurangi tingkat akurasi
3. Prosesor dalam perangkat *mobile* klien yang memiliki kemampuan komputasi yang rendah sehingga perlu dipertimbangkan metode yang tepat

Oleh karena itu, dibutuhkan beberapa strategi untuk mengurangi waktu komputasi yang dibutuhkan dan meningkatkan efisiensi model. Pada penelitian **(Konečný, dkk, 2016)** dan **(Zhu, dkk, 2019)**, proses agregasi dilakukan di server menggunakan metode *Federated Averaging* (FedAvg). Pada metode FedAvg, keseluruhan bobot yang diterima oleh server akan dilakukan perhitungan rata-rata bobot sesuai dengan banyaknya data yang diterima. Beberapa penelitian melakukan pengembangan proses agregasi menggunakan algoritma *swarm intelligence*, seperti yang dilakukan oleh **(Park, dkk, 2021) (Abasi, dkk, 2022) (Elfaki, dkk, 2023)**.

Model *federated learning* berbasis *particle swarm optimization* pada **(Park, dkk, 2021)** serta algoritma *grey wolf optimizer* **(Abasi, dkk, 2022)** bertujuan mengoptimasi bobot setiap klien. Pada setiap iterasi, hanya model terbaik saja yang dikirim ke server dan dijadikan sebagai

model global, sehingga tidak terdapat skema agregasi rata-rata. Salah satu kekurangan dari metode tersebut adalah setiap klien harus melakukan proses pembelajaran dengan *convolutional neural network* sekaligus optimasi bobot menggunakan algoritma optimasi pada setiap iterasi. Hal ini akan semakin meningkatkan beban komputasi dari setiap klien.

Dengan demikian, penelitian ini mengusulkan FedFA (*federated learning-firefly algorithm*) sebagai optimasi *communication cost* dengan metode optimasi bobot yang dilakukan hanya pada saat inisialisasi menggunakan algoritma *swarm intelligence* yang terinspirasi dari perilaku kunang-kunang, yakni algoritma *firefly*. Optimasi bobot inisialisasi dengan algoritma *firefly* telah dilakukan oleh **(Elakkiya E & Selvakumar S, 2019)** pada algoritma *feed forward neural network*. Pada penelitian ini, diasumsikan terdapat satu server global dan lima klien. Evaluasi kinerja FedFA diperoleh dari perhitungan akurasi dan *communication cost* berupa waktu komputasi, serta melakukan perbandingan kinerja FedFA dengan model orisinal FedAvg. Pengujian juga dilakukan dengan skema kondisi jaringan yang tidak stabil, sehingga data yang dikirim oleh klien tidak terkirim ke server. Dengan demikian, dapat diketahui pengaruhnya terhadap optimasi model.

2. METODE

Metode yang diusulkan dalam penelitian ini adalah FedFA: *Federated Learning with Firefly Algorithm* sebagai optimasi *communication cost* dalam proses komunikasi antara server dan klien pada komunikasi jaringan nirkabel. Metode yang terdapat dalam penelitian ini terdiri dari dua bagian yakni membuat model *federated learning* dan algoritma *firefly* sehingga dapat diintegrasikan menjadi model FedFA.

2.1 Model Federated Learning

Model *federated learning* terdiri dari *federated server* dan sejumlah klien lokal yang berkolaborasi dalam melatih sebuah model tanpa harus mengirimkan data mentah lokal ke server. Tahapan dalam model Federated Learning adalah sebagai berikut **(McMahan, dkk, 2017)**.

1. Update model lokal: setiap klien memperbarui model pembelajarannya secara lokal sesuai dengan data mentah yang dimiliki oleh setiap klien
2. Unggah nilai bobot: setiap klien mengirimkan parameter yang diperbarui dari model yang telah dilatih ke server
3. Agregasi: server melakukan agregasi parameter dari seluruh klien yang berpartisipasi
4. Feedback: server mengirimkan kembali parameter yang diperbarui ke setiap klien untuk iterasi berikutnya

Proses evaluasi kinerja dari model terbaru dilakukan dengan menghitung nilai dari *loss function*.

Algoritma 1: Federated Learning Framework

- 1: **Input:** Set klien $\mathcal{N} = \{1, 2, \dots, N\}$, iterasi lokal maksimum L , *learning rate* ζ , target akurasi A
- 2: **Output:** Global model w
- 3: **Inisialisasi** model global w^k , dengan $k = 1$
- 4: **for** *akurasi* $< A$
- 5: Server mengirimkan model global w^k ke seluruh klien pada set $\mathcal{N}$
- 6: **for**
- 7: setiap klien secara paralel lakukan
- 8: inisialisasi $i = 0$
- 9: **while** $i < L$ **do**

Algoritma 1: Federated Learning Framework

```

10:      setiap klien melakukan training data dan update lokal model
11:       $w_{n,k}(i+1) = w_{n,k}(i) - \zeta \nabla F_n(w_{n,k}(i));$ 
12:       $i = i + 1$ 
13:    end while
14:    setiap klien mengirimkan  $w_{n,k}(i+1)$  ke server
15:     $w_{n,k}(i+1) = w_{n,k}(i) - \zeta \nabla F_n(w_{n,k}(i))$ 
16:     $i = i + 1$ 
17:  end while
18:  setiap klien mengirimkan  $w_{n,k}(i+1)$  ke server
19: end for
20: Server mengagregasi seluruh bobot dari klien
21:  $w^k = \frac{\sum_{n \in S_k} D_n w_{n,k}}{\sum_{n \in S_k} D_n}$ 
22:    $k = k + 1$ 
23: Hitung akurasi
24: End

```

2.2 Firefly Algorithm

Algoritma firefly (Yang & He, 2013) adalah salah satu algoritma metaheuristik yang terinspirasi dari sifat hewan kunang-kunang. Terdapat sekitar 2000 spesies kunang-kunang dan sebagian besar kunang-kunang menghasilkan kilatan cahaya yang pendek dan berirama melalui bioluminesens. Cahaya yang berasal dari kunang-kunang adalah salah satu sistem sinyal untuk berkomunikasi dan berburu. Irama, durasi, dan tingkat kecerahan kilatan cahaya dapat bervariasi untuk setiap spesies yang berbeda, sehingga perilaku setiap spesies dapat diatur oleh masing-masing individu. Beberapa karakteristik khusus dari kunang-kunang adalah sebagai berikut:

- Kunang-kunang adalah hewan unisex, sehingga jenis kelamin kunang-kunang tidak diperhitungkan di dalam algoritma *firefly*
- Ketertarikan antar kunang-kunang berdasarkan tingkat kecerahan dari setiap kunang-kunang. Kunang-kunang yang lebih redup akan tertarik dan bergerak berpindah mendekati kunang-kunang yang lebih terang

Dalam algoritma *firefly*, dua parameter yang digunakan sebagai kendali untuk mencapai nilai optimal adalah variasi intensitas cahaya dan tingkat ketertarikan. Jika terdapat seekor kunang-kunang i pada posisi x_i dan seekor kunang-kunang j pada posisi x_j dengan jarak diantara keduanya adalah r_{ij} , maka

$$r_{ij} = \|x_i - x_j\| = \sqrt{\sum_{k=1}^d (x_{i,k} - x_{j,k})^2} \quad (1)$$

dengan $x_{i,k}$ adalah kunang-kunang x_i pada dimensi spasial ke- k .

Pergerakan dari kunang-kunang i mengikuti kunang-kunang lain yang memiliki tingkat kecerahan lebih terang yakni kunang-kunang j . Representasi persamaan matematis dari perpindahan kunang-kunang x_i adalah sebagai berikut.

$$x_i = x_i + \beta_0 e^{-\gamma r_{ij}^2} (x_j - x_i) + \alpha \epsilon_a \quad (2)$$

dengan α bernilai random dan ϵ_α adalah sebuah vektor dengan bilangan random yang terdistribusi normal. Dalam prakteknya, dapat digunakan nilai $\beta_0 = 1$ dan $\alpha \in [0,1]$. Koefisien penyerapan cahaya γ mempengaruhi kecepatan konvergensi dan karakteristik algoritma, pada umumnya bernilai antara 0,1 - 10. Persamaan yang digunakan untuk memperbarui nilai x_i adalah metode *random walk* dari seekor kunang-kunang menuju kunang-kunang yang lebih terang.

2.3 Integrasi *Firefly Algorithm* dengan *Federated Learning*

Pada penelitian ini, algoritma *firefly* digunakan untuk mengoptimasi bobot inisialisasi pada setiap klien. Prinsip dasar dari FedFA disajikan pada Algoritma 2. Inisialisasi sebuah jaringan yang terdiri dari satu server dan satu set klien $\mathcal{N} = \{1, 2, \dots, N\}$, *client epoch* C , *learning rate* ζ , dan iterasi global G . Server mengirimkan model global ke seluruh klien. Setiap klien merupakan representasi dari *firefly*. Model klien yang memiliki nilai *loss* paling rendah dijadikan model global. Setelah optimasi bobot inisialisasi, maka prosedur update model global selanjutnya dilakukan menggunakan metode FedAvg.

Algoritma 2: FedFA (*Federated Learning with Firefly Algorithm*)

```

1: Input: Set klien  $\mathcal{N} = \{1, 2, \dots, N\}$ , client epoch  $C$ , learning rate  $\zeta$ , iterasi optimasi  $G$ , iterasi global  $I$ 
2: Output: Global model  $w$ 
3: Inisialisasi model global  $w^k$ , dengan  $k = 1$ 
4: Server mengirim model global  $w^k$  ke seluruh klien pada set  $\mathcal{N}$ , lakukan prosedur optimasi inisialisasi bobot:
   while (iterasi <  $G$ )
5:   for setiap klien  $i$  secara paralel lakukan
6:     Inisialisasi koefisien penyerapan cahaya  $\gamma$ 
7:     Hitung akurasi dan loss pada setiap klien berdasarkan dataset masing-masing klien.
8:     for  $j = 1 : n$  untuk sejumlah  $n$  kunang-kunang
9:       if  $i \neq j$ 
10:        if ( $I_i < I_j$ ), pindahkan kunang-kunang  $i$  menuju  $j$ ; end if
11:        untuk setiap layer, lakukan:
12:          Hitung jarak Euclidian  $r$  pada nilai bobot kunang-kunang  $i$  dan kunang-kunang  $j$ 
13:          Variasikan tingkat ketertarikan terhadap jarak dengan  $e^{-\gamma r_{i,j}^2}$ 
14:          Update posisi terbaru kunang-kunang  $i$ 
15:          while (epoch <  $C$ )
16:            Setiap klien melakukan training data dan update lokal model dengan bobot baru
17:             $w_{n,k}(t+1) = w_{n,k}(t) - \zeta \nabla F_n(w_{n,k}(t))$ ;
18:            Simpan nilai loss ( $l_i$ )
19:          Setiap klien mengirimkan nilai loss ke Server
20:          Server mengurutkan nilai loss
21:          Klien dengan nilai loss terkecil mengirimkan lokal model ke server
22:        End
23: Server mengirim model global  $w^k$  ke seluruh klien pada set  $\mathcal{N}$ , lakukan prosedur FedAvg pada Algoritma 1

```

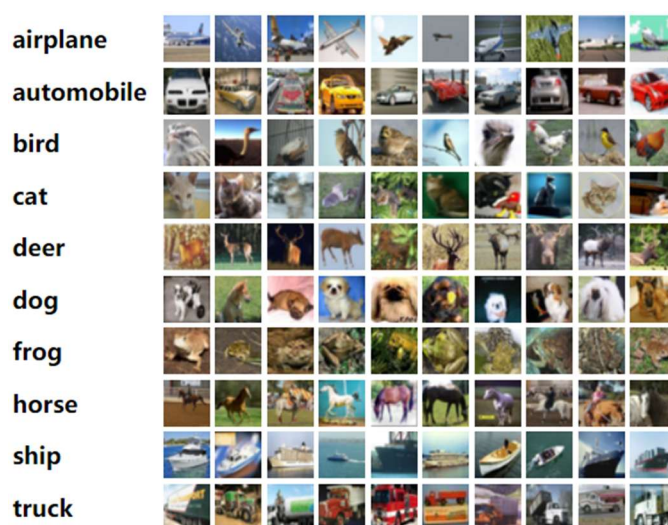
3. HASIL DAN PEMBAHASAN

3.1 Setup Eksperimen

Evaluasi dari kinerja algoritma FedFA terdiri dari evaluasi kecepatan konvergensi dan akurasi. Kinerja FedFA dibandingkan dengan algoritma orisinal *federated learning* yakni FedAvg. Simulasi dilakukan menggunakan perangkat keras dengan spesifikasi prosesor i7 13th Gen Intel(R) Core(TM) 2,6 GHz, NVIDIA GeForce RTX 4050. Model pembelajaran yang dilakukan menggunakan algoritma *convolutional neural network* (CNN), identik dengan algoritma orisinal dari *federated learning* yakni FedAvg (McMahan, dkk, 2017).

Dataset yang digunakan adalah *Canadian Institute for Advanced Research's* (CIFAR-10) (Alex Krizhevsky dkk., 2009). Dataset CIFAR-10 (Gambar 1) terdiri dari 60000 32x32 citra berwarna yang terbagi menjadi 10 kelas (pesawat udara, mobil, burung, kucing, rusa, anjing, katak, kuda, kapal, dan truk), masing-masing kelas berjumlah 6000 citra. Terdapat 50000 citra untuk proses pembelajaran serta 10000 citra untuk proses pengujian. Pada proses pembelajaran *federated learning*, diasumsikan seluruh klien memiliki jenis dataset IID (*independent and identically distributed*).

Pada saat proses pembelajaran, dataset CIFAR-10 dialokasikan kepada seluruh klien dengan jumlah yang sama dengan metode pembagian data yang acak. Arsitektur CNN yang digunakan terlihat pada Tabel 1, sedangkan parameter algoritma FedFA dicantumkan pada Tabel 2.



Gambar 1. Dataset CIFAR-10 (Alex Krizhevsky dkk, 2009)

Tabel 1. Parameter CNN

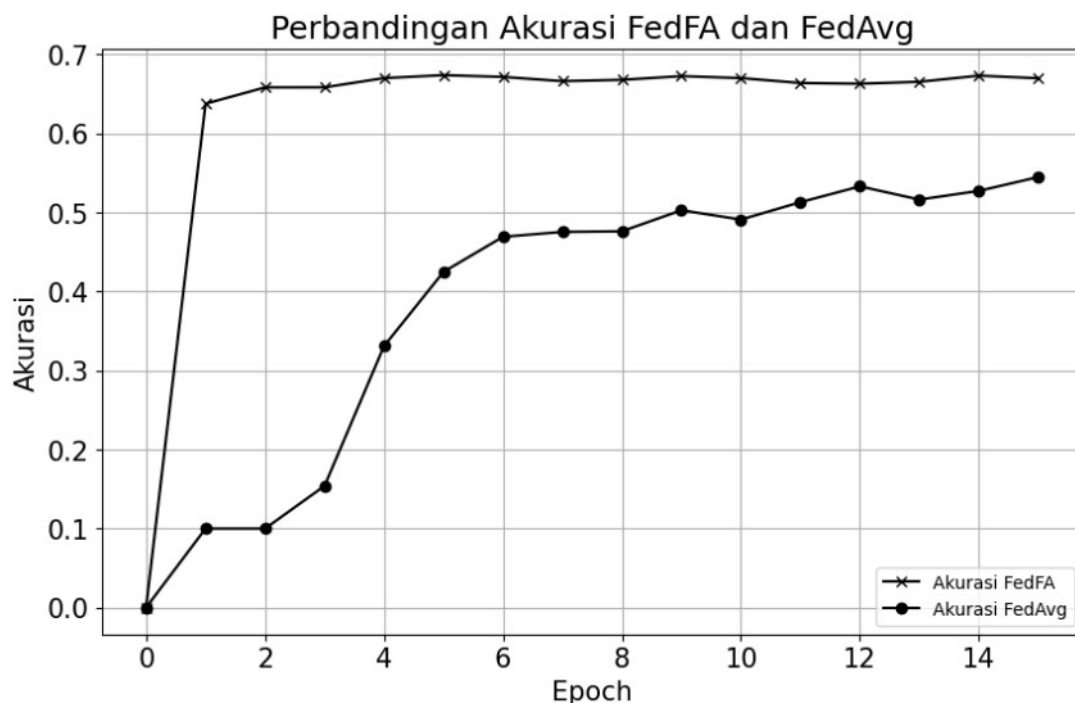
No	Shape	Layer
1	5 x 5 x 32	Conv2D
2	32	Conv2D
3	5 x 5 x 64	Conv2D
4	64	Conv2D
5	1024 x 512	Dense
6	512	Dense
7	512 x 10	Dense
8	10	Dense

Tabel 2. Parameter algoritma FedFA

No	Parameter	FedFA	FedAvg
1	Client-epoch	5	5
2	<i>Communication round (server-client epoch)</i>	15	15
3	Batch	10	10
4	Jumlah client	5	5

3.2 Pengujian Nilai Akurasi

Gambar 2 menunjukkan nilai akurasi dari pengujian menggunakan data uji dari dataset CIFAR-10. Hasil pengujian menunjukkan peningkatan akurasi menjadi 67,31% jika dibandingkan dengan FedAvg yang hanya mencapai 54,47% hingga iterasi ke-15. Berdasarkan hasil tersebut, dapat disimpulkan bahwa optimasi inisialisasi bobot menggunakan algoritma *firefly* memiliki dampak yang cukup signifikan dalam meningkatkan akurasi pada model global.

**Gambar 2. Perbandingan akurasi FedFA dan FedAvg**

3.3 Pengujian Waktu Komputasi

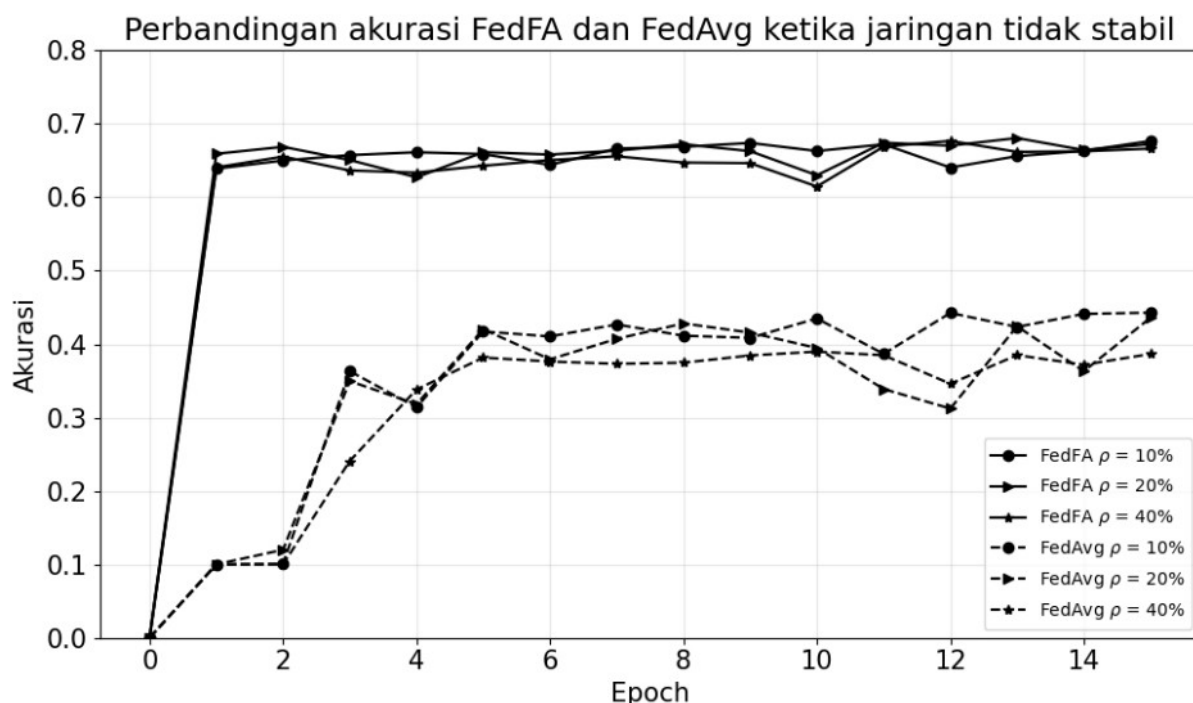
Algoritma yang diusulkan, FedFA, menjalankan optimasi bobot inisialisasi terlebih dahulu. Hal ini dilakukan untuk mengurangi waktu yang dibutuhkan untuk menentukan model global selama proses pembelajaran. Pengujian waktu komputasi dilakukan dengan membandingkan algoritma FedFA dan algoritma orisinal FedAvg dengan perangkat keras yang sama. Hasil pengujian waktu komputasi disajikan dalam Tabel 3. Waktu yang dibutuhkan untuk mencapai target akurasi 50% pada algoritma FedFA adalah 158,33 detik, sedangkan algoritma FedAvg membutuhkan waktu yang lebih lama yakni 992,49 detik. Untuk mencapai target akurasi 60%, algoritma FedFA adalah 807,92 detik. Dengan demikian, dapat disimpulkan bahwa meskipun terdapat optimasi bobot pada saat inisialisasi, waktu komputasi yang diperlukan untuk mencapai target akurasi lebih singkat pada algoritma FedFA.

Tabel 3. Perbandingan waktu komputasi untuk mencapai nilai akurasi 50% dan 60%

Model	Waktu Komputasi	
	Akurasi 50%	Akurasi 60%
FedFA	158,33 detik	807,92 detik
FedAvg	992,49 detik	-

3.4 Pengujian saat Jaringan Komunikasi Tidak Stabil

Federated learning menggunakan komunikasi nirkabel sebagai media transmisi data. Komunikasi nirkabel memiliki beberapa karakteristik, salah satunya adalah jaringan komunikasi yang tidak selalu stabil. Oleh karena itu, pengujian berikut mengasumsikan beberapa klien tidak dapat mengirimkan bobot ke server untuk menguji tingkat akurasi pada kondisi jaringan yang tidak stabil. Tingkat kegagalan komunikasi diatur berdasarkan parameter probabilitas kegagalan komunikasi (ρ_f). Parameter nilai akurasi dibandingkan berdasarkan variasi nilai ρ_f yakni 10%, 20%, dan 40%.

**Gambar 3. Perbandingan akurasi FedFA dan FedAvg ketika jaringan tidak stabil**

Berdasarkan hasil pengujian yang disajikan pada Gambar 3, terlihat bahwa algoritma FedFA memiliki akurasi yang lebih baik yakni 67,67% pada saat kegagalan 10%, 67,24% saat kegagalan komunikasi 20%, 66,63% saat kegagalan komunikasi 40%. Dengan kata lain, ρ_f senilai 10% dan 20% tidak terlalu berpengaruh terhadap tingkat akurasi, sedangkan kegagalan komunikasi 40% menurunkan akurasi sekitar 1%. Di sisi lain, algoritma FedAvg mengalami penurunan akurasi yang signifikan yakni akurasi sebesar 44,29% saat kegagalan komunikasi 10%, akurasi 43,6% saat kegagalan komunikasi 20%, dan akurasi 38,66% saat kegagalan komunikasi 40%. Hasil tersebut menunjukkan bahwa model global dari algoritma FedFA lebih adaptif dalam proses pembelajaran meskipun terdapat klien yang tidak dapat mengirimkan bobot ke server dan tidak terlibat dalam proses pembelajaran pada suatu *communication round*.

4. KESIMPULAN

Penelitian ini mengusulkan algoritma FedFA berbasis algoritma *firefly* untuk meningkatkan kinerja komunikasi jaringan dari *federated learning* dengan mempersingkat waktu komunikasi, sehingga *communication cost* lebih optimal. Inisialisasi bobot dilakukan menggunakan algoritma *firefly*. Algoritma yang diusulkan menghasilkan peningkatan akurasi sebesar 12,84% dibandingkan dengan FedAvg. Model FedFA juga lebih tahan terhadap gangguan komunikasi, terlihat dari penurunan akurasi yang tidak terlalu signifikan jika dibandingkan dengan algoritma FedAvg.

UCAPAN TERIMA KASIH

Ucapan terimakasih diberikan kepada Direktorat Jenderal Pendidikan Tinggi, Riset, dan Teknologi Kementerian Pendidikan, Kebudayaan, Riset, dan Teknologi atas pendanaan hibah riset skema Penelitian Dosen Pemula Tahun Pendanaan 2024 dengan Nomor Kontrak Penelitian 105/E5/PG.02.00.PL/2024, 788/LL3/AL.04/2024, A/UBL/DRPM/000/062/06/24, penelitian ini dapat dilaksanakan dan dipublikasikan pada Jurnal ELKOMIKA.

DAFTAR RUJUKAN

- Abasi, A. K., Aloqaily, M., & Guizani, M. (2022). Grey Wolf Optimizer for Reducing Communication Cost of Federated Learning. *Proceedings - IEEE Global Communications Conference, GLOBECOM*, (pp. 1049–1054). <https://doi.org/10.1109/GLOBECOM48099.2022.10001681>
- Alex Krizhevsky, Vinod Nair, & Geoffrey Hinton. (n.d.). *The CIFAR-10 Dataset*. <https://www.cs.toronto.edu/~kriz/cifar.html>.
- Brendan McMahan Eider Moore Daniel Ramage Seth Hampson Blaise Agüera-Ag, H., & Arcas, A. (2017). *Communication-Efficient Learning of Deep Networks from Decentralized Data*.
- Cisco. (2020). Cisco Annual Internet Report (2018-2023). *White Paper Cisco Public, 2020*(3), 4–4. [https://doi.org/10.1016/s1361-3723\(20\)30026-9](https://doi.org/10.1016/s1361-3723(20)30026-9)
- Elakkiya E, & Selvakumar S. (2019). Initial Weights Optimization using Enhanced Step Size Firefly Algorithm for Feed Forward Neural Network applied to Spam Detection. *TENCON 2019 - 2019 IEEE Region 10 Conference (TENCON): Proceedings: Technology, Knowledge, and Society*, 1362.
- Elfaki, M. A., Alshahrani, H. M., Mahmood, K., Alabdan, R., Alymani, M., Alshahrani, H., Motwakel, A., & Alneil, A. A. (2023). Metaheuristics Algorithm-Based Minimization of Communication Costs in Federated Learning. *IEEE Access*, 11, 81310–81317. <https://doi.org/10.1109/ACCESS.2023.3300221>

- Konečný, J., McMahan, H. B., Yu, F. X., Richtárik, P., Suresh, A. T., & Bacon, D. (2016). *Federated Learning: Strategies for Improving Communication Efficiency*. <http://arxiv.org/abs/1610.05492>
- Li, T., Sahu, A. K., Talwalkar, A., & Smith, V. (2020). Federated Learning: Challenges, Methods, and Future Directions. *IEEE Signal Processing Magazine*, 37(3), 50–60. <https://doi.org/10.1109/MSP.2020.2975749>
- Lu, Y., Huang, X., Zhang, K., Maharjan, S., & Zhang, Y. (2021). Blockchain and Federated Learning for 5G beyond. *IEEE Network*, 35(1), 219–225. <https://doi.org/10.1109/MNET.011.1900598>
- Nguyen, D. C., Ding, M., Pathirana, P. N., Seneviratne, A., Li, J., & Poor, H. V. (2021). Federated Learning for Internet of Things: A Comprehensive Survey. *IEEE Communications Surveys & Tutorials*, 23(3). <https://doi.org/10.1109/COMST.2021.3075439>
- Niknam, S., Dhillon, H. S., & Reed, J. H. (2020). Federated Learning for Wireless Communications: Motivation, Opportunities, and Challenges. *IEEE Communications Magazine*, 58(6), 46–51. <https://doi.org/10.1109/MCOM.001.1900461>
- Park, S., Suh, Y., & Lee, J. (2021). FedPSO: Federated learning using particle swarm optimization to reduce communication costs. *Sensors (Switzerland)*, 21(2), 1–13. <https://doi.org/10.3390/s21020600>
- Qin, Z., Li, G. Y., & Ye, H. (2021). Federated Learning and Wireless Communications. *IEEE Wireless Communications*, 28(5), 134–140. <https://doi.org/10.1109/MWC.011.2000501>
- Yang, X. S., & He, X. (2013). Firefly algorithm: recent advances and applications. *International Journal of Swarm Intelligence*, 1(1), 36. <https://doi.org/10.1504/ijsi.2013.055801>
- Zhang, T., Gao, L., He, C., Zhang, M., Krishnamachari, B., & Avestimehr, A. S. (2022). Federated Learning for the Internet of Things: Applications, Challenges, and Opportunities. *IEEE Internet of Things Magazine*, 5(1), 24–29. <https://doi.org/10.1109/iotm.004.2100182>
- Zhu, L., Liu, Z., & Han, S. (2019). Deep Leakage from Gradients. *Advances in Neural Information Processing Systems 32 (NeurIPS 2019)*. <http://arxiv.org/abs/1906.08935>