



Artikel

Memahami Sentimen Pengguna: Analisis Ulasan Aplikasi STRAVA di *Google Play Store*

Muhit Fahri Alamsah ^{1,*} dan Rian Ardianto ²

¹ Departemen Sistem Informasi, Universitas Harapan Bangsa, Purwokerto, Indonesia

² Departemen Informatika, Universitas Harapan Bangsa, Purwokerto, Indonesia

*Korespondensi: muhidfahrialamsah@gmail.com

Abstrak: Penelitian ini bertujuan untuk menganalisis sentimen pengguna terhadap aplikasi STRAVA berdasarkan ulasan yang diperoleh dari *Google Play Store*. Data yang digunakan berjumlah 6.082 ulasan setelah melalui tahap pembersihan dan pra-pemrosesan teks, termasuk *case folding*, *stopword removal*, *tokenizing*, dan *stemming*. Ulasan dengan *rating* 4–5 dikategorikan sebagai sentimen positif, sedangkan *rating* 1–2 dikategorikan sebagai sentimen negatif. Algoritma *Naive Bayes* digunakan sebagai model klasifikasi dengan representasi fitur berbasis *Term Frequency-Inverse Document Frequency* (TF-IDF). Hasil evaluasi menunjukkan bahwa model mampu mengklasifikasikan sentimen dengan akurasi sebesar 90,7%, presisi dan *recall* pada sentimen positif mencapai 94%, sementara pada sentimen negatif masing-masing sekitar 78%. Temuan ini mengindikasikan bahwa *Naive Bayes* dengan pembobotan TF-IDF efektif dalam menganalisis sentimen ulasan aplikasi STRAVA. Penelitian ini diharapkan dapat memberikan masukan bagi pengembang STRAVA dalam meningkatkan kualitas fitur, pengalaman pengguna, serta kepuasan secara keseluruhan.

Received: 12 November 2024

Revised: 30 Desember 2024

Accepted: 30 Januari 2025

Published: 6 Mei 2025



Copyright: © 2023 by the authors.

License Universitas Harapan Bangsa, Purwokerto, Indonesia. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

Kata kunci: Analisis Sentimen; *Text Mining*; *Naive Bayes*; TF-IDF; STRAVA

Pendahuluan

Olahraga dan aktivitas fisik merupakan bagian penting dari gaya hidup sehat masyarakat modern (Mulyana et al., 2024; Sopa & Pomohaci, 2018). Perkembangan teknologi digital telah mendorong hadirnya berbagai aplikasi kebugaran berbasis *mobile*, salah satunya STRAVA (Heidianto & Wijayanti, 2025; Hochstetler, 2024). Aplikasi ini digunakan oleh jutaan pengguna di seluruh dunia untuk merekam, memantau, dan berbagi aktivitas olahraga seperti berlari, bersepeda, berenang, dan *hiking* (Hochstetler, 2024; Robinson et al., 2024). STRAVA memanfaatkan teknologi GPS untuk merekam rute dan performa pengguna, sekaligus menyediakan fitur

komunitas yang memungkinkan interaksi serta motivasi sosial (Nam & Kwon, 2024). Seiring meningkatnya jumlah pengguna, aplikasi STRAVA menerima berbagai ulasan di *Google Play Store* yang berisi tanggapan, kritik, maupun pujian terkait pengalaman pengguna (Heidianto & Wijayanti, 2025). Ulasan-ulasan tersebut menjadi sumber informasi yang berharga bagi pengembang untuk memahami tingkat kepuasan dan ekspektasi pengguna. Namun, jumlah ulasan yang sangat banyak membuat proses analisis secara manual menjadi tidak efisien. Oleh karena itu, diperlukan pendekatan analisis sentimen untuk mengolah data teks ulasan secara otomatis dan mengklasifikasikannya ke dalam kategori positif, negatif, atau netral (Rivers, 2020).

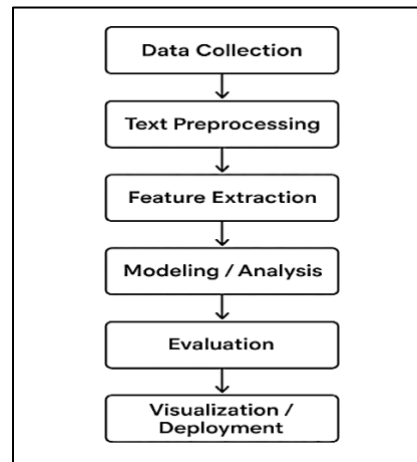
Analisis sentimen sendiri merupakan bagian dari *opinion mining* yang memanfaatkan teknik *text mining* untuk mengekstrak informasi subjektif dari teks (Sun, 2017). Salah satu metode yang populer dalam analisis sentimen adalah *Naive Bayes Classifier* (NBC), yang dikenal memiliki keunggulan dari segi kesederhanaan, kecepatan, serta akurasi yang tinggi dalam pengolahan teks (Byun et al., 2023). Berbagai penelitian menunjukkan bahwa NBC unggul dibandingkan metode lain seperti *Support Vector Machine* atau *Decision Tree*, terutama pada dataset berbahasa alami dengan jumlah data besar. Penerapan NBC pada analisis ulasan aplikasi di *Google Play Store*, seperti pada aplikasi *Tije*, *Shopee*, *Gojek*, dan *KitaLulus*, menghasilkan tingkat akurasi yang berkisar antara 80% hingga 91%. Kombinasi NBC dengan pembobotan kata menggunakan *Term Frequency-Inverse Document Frequency* (TF-IDF) terbukti mampu meningkatkan kinerja model dalam mengklasifikasikan *sentiment* (Di Sorbo et al., 2021). Selain itu, tahap *pre-processing* seperti penghapusan kata-kata tidak penting (*stopword removal*), normalisasi huruf (*case folding*), dan *stemming* menjadi langkah krusial untuk mencapai akurasi yang optimal.

Sebagai aplikasi populer di bidang olahraga, kemunculan STRAVA memunculkan beragam reaksi dari para penggunanya, mulai dari apresiasi terhadap fitur inovatif hingga keluhan mengenai *bug* dan masalah teknis. Perbedaan opini ini membuat analisis sentimen menjadi penting untuk memahami persepsi publik secara lebih sistematis. Hasil analisis dapat dimanfaatkan oleh pengembang untuk melakukan perbaikan fitur, meningkatkan performa aplikasi, serta mendorong kepuasan pengguna (Venter et al., 2023). Dengan latar belakang tersebut, penelitian ini dilakukan untuk mengklasifikasikan ulasan pengguna STRAVA di *Google Play Store* menjadi sentimen positif dan negatif, mengevaluasi kinerja algoritma *Naive Bayes* dengan pembobotan TF-IDF pada dataset ulasan, serta memberikan rekomendasi berbasis data bagi pengembang STRAVA guna meningkatkan kualitas layanan. Penelitian ini diharapkan dapat memberikan kontribusi bagi pengembangan aplikasi STRAVA sekaligus menjadi referensi penerapan *machine learning* dalam analisis sentimen pada aplikasi berbasis komunitas olahraga.

Penelitian ini memiliki nilai kebaruan karena mengangkat konteks aplikasi STRAVA, sebuah aplikasi kebugaran berbasis komunitas yang belum banyak dianalisis sentimennya di Indonesia melalui pendekatan *supervised learning*. Sebagian besar penelitian sebelumnya berfokus pada *e-commerce* atau layanan transportasi daring, sementara studi terhadap aplikasi kebugaran masih terbatas. Selain itu, penelitian ini menggabungkan pendekatan analisis ulasan pengguna *Google Play Store* dengan metode klasifikasi berbasis *Naive Bayes* dan pembobotan TF-IDF yang divalidasi secara eksperimental melalui pengujian akurasi, presisi, dan *recall*. Penelitian ini juga memberikan gambaran kontribusi praktis terhadap pengembangan aplikasi berbasis komunitas dengan wawasan berbasis data.

Metode

Penelitian ini mengadopsi pendekatan CRISP-DM (*Cross-Industry Standard for Data Mining*) seperti ditunjukkan pada Gambar 1. Tahapan CRISP-DM meliputi *Business Understanding*, *Data Understanding*, *Data Preparation*, *Modeling*, *Evaluation*, dan *Deployment* (Chapman, 1999; Rianti et al., 2023).



Gambar 1. Diagram alur penelitian

Business Understanding

Tahap pemahaman bisnis berfokus pada pengenalan objek penelitian. Pada penelitian ini, objek yang dikaji adalah aplikasi STRAVA, sebuah platform pelacak aktivitas kebugaran yang digunakan untuk merekam dan berbagi aktivitas seperti berlari, bersepeda, dan berenang. Pemahaman objek penelitian dilakukan dengan mengumpulkan data ulasan pengguna dari *Google Play Store* menggunakan alat *web scraper* pada *Google Colab*, khususnya yang menargetkan ulasan aplikasi STRAVA. Langkah ini didasari oleh fakta bahwa ulasan pengguna memuat informasi tekstual yang mencerminkan pengalaman, opini, dan masukan terkait kinerja aplikasi. Platform ulasan daring tidak hanya menjadi wadah opini, tetapi juga dapat digunakan untuk memantau tren sentimen publik terhadap suatu produk atau layanan. Analisis sentimen digunakan untuk menemukan pendekatan klasifikasi yang dapat mengidentifikasi sentimen positif dan negatif dari ulasan pengguna. Pada tahap ini juga dibentuk pemahaman mengenai algoritma klasifikasi yang paling sesuai untuk mendukung proses pengolahan data berikutnya.

Data Understanding

Proses memperoleh data mentah berdasarkan atribut yang relevan dilakukan pada tahap pemahaman data. Data dikumpulkan dari *Google Play Store* pada tautan berikut:

["https://play.google.com/store/apps/details?id=com.strava&hl=en-ID"](https://play.google.com/store/apps/details?id=com.strava&hl=en-ID)

Informasi diambil dari ulasan dengan penilaian (*rating*) antara 1 hingga 5 bintang. Total data mentah yang terkumpul berjumlah 6.843 ulasan (d disesuaikan dengan hasil pengumpulan data). Setelah proses pengumpulan data, dilakukan pembersihan data sehingga diperoleh 5.124 ulasan yang dapat digunakan. Penelitian ini menggunakan data ulasan dalam bahasa Indonesia dan bahasa Inggris, dengan langkah pra-pemrosesan yang memastikan konsistensi penanganan bahasa.

Data ulasan aplikasi STRAVA diambil dari platform *Google Play Store* menggunakan teknik *web scraping*. Proses pengambilan data dilakukan selama lima hari, yaitu pada tanggal 12 hingga 16 Juni 2023, dengan menggunakan pustaka *google-play-scraper* dalam bahasa pemrograman *Python*. Jumlah ulasan yang dikumpulkan adalah 5.000 komentar pengguna dengan batasan waktu 1 tahun terakhir dari tanggal pengambilan data. Seluruh data yang

diperoleh masih berupa teks mentah yang bercampur antara Bahasa Indonesia dan Bahasa Inggris. Untuk menjaga kesesuaian konteks, data yang digunakan dalam proses pelabelan hanya dipilih yang menggunakan Bahasa Indonesia, sementara ulasan berbahasa Inggris atau yang menggunakan simbol dan ekspresi emosional tanpa kata (misal: emoji saja) dikeluarkan dari dataset melalui proses *data cleaning* awal.

Data Preparation

Tahap persiapan data bertujuan menghasilkan data yang bersih dan siap digunakan untuk keperluan penelitian. Langkah *pre-processing* teks dilakukan pada tahap awal *text mining* menggunakan alat di *Google Colab*. Beberapa teknik *pre-processing* yang diterapkan pada dataset meliputi:

- *Filtering*: menghapus entri duplikat, ulasan non-teks, dan komentar yang terlalu singkat.
- *Labeling*: memetakan ulasan dengan *rating* 4–5 sebagai sentimen positif, *rating* 1–2 sebagai sentimen negatif, dan mengabaikan ulasan netral dengan *rating* 3.
- *Cleaning & Case Folding*: mengubah semua teks menjadi huruf kecil, menghapus tanda baca, karakter khusus, dan spasi berlebih.
- *Stopword Removal*: menghapus kata-kata umum yang tidak memiliki nilai informasi signifikan dalam bahasa Indonesia dan bahasa Inggris.
- *Tokenizing*: memecah kalimat menjadi token kata.
- *Stemming*: mengembalikan kata ke bentuk dasarnya untuk menyatukan variasi kata yang sama.

Modeling

Tahap ini menentukan teknik klasifikasi yang akan digunakan. Penelitian ini menggunakan algoritma *Naive Bayes* sebagai model klasifikasi. Proses pemodelan dilakukan di *Google Colab* menggunakan pustaka *Scikit-Learn*. Dataset yang telah melalui proses *pre-processing* diubah menjadi representasi numerik menggunakan *Term Frequency-Inverse Document Frequency* (TF-IDF) sebelum dimasukkan ke *Naive Bayes Classifier*. Model dilatih untuk mengklasifikasikan ulasan menjadi kelas sentimen positif dan negatif.

Evaluation

Tahap evaluasi bertujuan menilai efektivitas model yang dikembangkan. Model klasifikasi sentimen pada ulasan STRAVA diuji menggunakan dataset yang dibagi menjadi 80% data latih dan 20% data uji. Metrik kinerja seperti akurasi, presisi, *recall*, dan *F1-score* dihitung dengan menggunakan *Confusion Matrix* sebagai acuan. Pembobotan kata dengan TF-IDF diterapkan pada proses ekstraksi fitur untuk meningkatkan akurasi klasifikasi.

Deployment

Tahap *deployment* adalah pengembangan model implementasi yang dapat diintegrasikan ke berbagai platform atau alat pemantauan. Hasil evaluasi pada penelitian ini dapat digunakan sebagai data acuan untuk pengembangan selanjutnya, sehingga memungkinkan pemantauan sentimen ulasan STRAVA secara otomatis. Hal ini dapat membantu tim pengembang STRAVA dalam mengidentifikasi permasalahan dengan cepat, melakukan perbaikan fitur, dan meningkatkan kepuasan pengguna.

Pembobotan Kata

Pembobotan kata adalah proses pemberian nilai pada setiap kata berdasarkan tingkat kepentingan dan pengaruhnya terhadap hasil klasifikasi. Nilai ini dapat digunakan sebagai dasar pemilihan fitur. Pendekatan yang digunakan adalah TF-IDF (*Term Frequency-Inverse Document Frequency*), salah satu metode yang banyak digunakan dalam *text mining*. Rumus TF-IDF dijabarkan dalam Persamaan (1).

$$W_{i,j} = \frac{n_{i,j}}{\sum_{j=1}^p n_{i,j}} \times \log_2 \frac{D}{d_j} \quad (1)$$

Penjelasan:

$W_{i,j}$ = bobot TF-IDF untuk term "j" pada dokumen "i"

$n_{i,j}$ = jumlah kemunculan term "j" pada dokumen "i"

p = total jumlah term pada dokumen

D = total jumlah dokumen

d_j = jumlah dokumen yang memuat term "j"

Pemrosesan Data

Algoritma *Naive Bayes* adalah metode klasifikasi berbasis statistik yang digunakan untuk menghitung probabilitas sebuah data masuk ke kelas tertentu. *Naive Bayes* dikenal memiliki kecepatan tinggi dan akurasi yang baik dalam memproses dataset teks berukuran besar [4]. Rumus umum *Naive Bayes* dijelaskan pada Persamaan (2).

$$P(H|X) = \frac{P(H) \times P(X|H)}{P(X)} \quad (2)$$

Keterangan :

X = data dengan kelas yang belum diketahui

H = hipotesis bahwa data X termasuk ke kelas tertentu

$P(H|X)$ = probabilitas hipotesis H dengan kondisi data X

$P(H)$ = probabilitas hipotesis H

$P(X|H)$ = probabilitas data X muncul dengan hipotesis H

Confusion Matrix digunakan untuk mengevaluasi prediksi model terhadap nilai *True Positive* (TP), *False Positive* (FP), *True Negative* (TN), dan *False Negative* (FN) (Al-Hakim & Prokopchuk, 2024). Nilai akurasi, presisi, dan *recall* dihitung dengan Persamaan (3).

$$Akurasi = \frac{TP + TN}{TP + TN + FP + FN}$$

$$Presisi = \frac{TP}{TP + FP} \times 100\%$$

$$Recall = \frac{TP}{TP + FN} \times 100\%$$

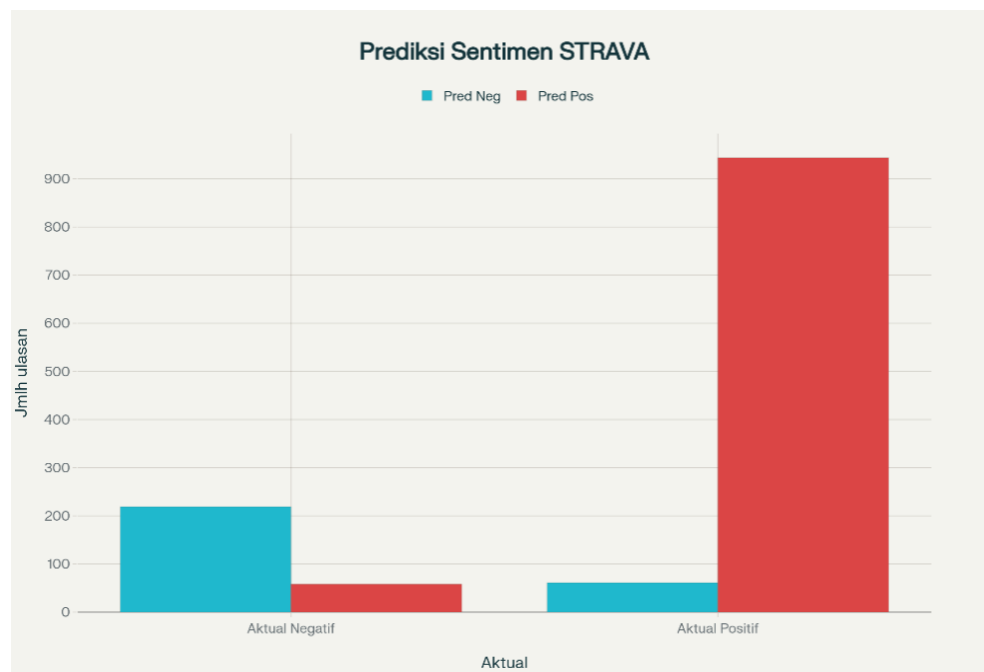
(3)

Hasil dan Pembahasan

Hasil

Penelitian ini menggunakan data ulasan aplikasi STRAVA yang diambil dari *Google Play Store* dengan jumlah total sekitar 6.000 ulasan. Setelah dilakukan proses pembersihan data (*data cleaning*), setiap ulasan dikategorikan menjadi kelas positif atau negatif berdasarkan skor ulasan. Dataset akhir yang digunakan untuk analisis sentimen menggunakan algoritma *Naive Bayes* dengan pembobotan TF-IDF berjumlah 6.082 ulasan setelah *filtering* dan pelabelan.

Tahap *pre-processing* meliputi *filtering* data, pelabelan, *case folding*, penghapusan *stopword*, *tokenizing*, dan *stemming* untuk menyiapkan data sebelum klasifikasi. Dataset dibagi menjadi 80% data latih (4.800 ulasan) dan 20% data uji (1.282 ulasan), dengan distribusi data uji 277 ulasan negatif dan 1.005 positif. Gambar 2 menjelaskan grafik hasil prediksi sentimen yang dilakukan.



Gambar 2. Grafik batang hasil prediksi sentimen

Pembahasan

Hasil pengujian menunjukkan bahwa algoritma *Naive Bayes* yang dipadukan dengan pembobotan TF-IDF mampu mengklasifikasikan sentimen ulasan aplikasi STRAVA dengan tingkat akurasi yang cukup tinggi, yaitu 90,7%. Hal ini menunjukkan bahwa model dapat secara efektif membedakan antara ulasan positif dan negatif.

Presisi dan *recall* yang tinggi pada kelas positif (94%) menandakan bahwa model sangat baik dalam mengidentifikasi ulasan yang bersifat positif. Sedangkan pada kelas negatif, nilai presisi dan *recall* masing-masing di kisaran 78%, yang masih tergolong cukup baik meskipun terdapat ketidakseimbangan jumlah ulasan (lebih banyak ulasan positif daripada negatif). Hal ini mengindikasikan tantangan yang umum pada data tidak seimbang dalam analisis sentimen.

Performa ini konsisten dengan penelitian sebelumnya yang juga menggunakan *Naive Bayes* untuk analisis sentimen ulasan aplikasi di *Google Play Store*, dengan akurasi yang berkisar antara 80% hingga 91%. Hasil penelitian ini menandakan aplikasi cukup baik diterima masyarakat sebagai alat gaya hidup baru, sejalan dengan penelitian lain yang serupa dilakukan oleh Kim dan Chung (2024).

Kesimpulan

Penelitian ini menunjukkan bahwa algoritma *Naive Bayes* dengan pembobotan TF-IDF mampu mengklasifikasikan ulasan pengguna aplikasi STRAVA di *Google Play Store* dengan tingkat akurasi yang tinggi, yaitu 90,7%. Model terbukti efektif dalam mengenali sentimen positif dengan presisi dan *recall* mencapai 94%, meskipun performa pada sentimen negatif masih lebih rendah akibat ketidakseimbangan data. Hasil ini menegaskan bahwa metode *text mining* berbasis *Naive Bayes* dapat digunakan secara efisien untuk memahami persepsi pengguna terhadap aplikasi kebugaran. Bagi pengembang STRAVA, hasil analisis ini dapat dimanfaatkan untuk mengidentifikasi aspek yang mendapat apresiasi maupun kritik, sehingga menjadi dasar dalam perbaikan fitur dan peningkatan kualitas layanan.

Referensi

- Al-Hakim, R. R., & Prokopchuk, Y. (2024). Predicting Internal Diseases in Humans Using Machine Learning: A Systematic Literature Review. *Journal of Advanced Health Informatics Research (JAHIR)*, 2(1), 50–63. <https://doi.org/10.59247/jahir.v2i1.195>
- Byun, H., Chiu, W., & Won, D. (2023). The Voice from Users of Running Applications: An Analysis of Online Reviews Using Leximancer. *Journal of Theoretical and Applied Electronic Commerce Research*, 18(1), 173–186. <https://doi.org/10.3390/jtaer18010010>
- Chapman, P. (1999). The CRISP-DM User Guide. *4th CRISP-DM SIG Workshop in Brussels*.
- Di Sorbo, A., Grano, G., Aaron Visaggio, C., & Panichella, S. (2021). Investigating the criticality of user-reported issues through their relations with app rating. *Journal of Software: Evolution and Process*, 33(3). <https://doi.org/10.1002/smr.2316>
- Heidianto, H. A., & Wijayanti, D. G. S. (2025). Motivation of Gen Z to use The Strava Application as A Way to Maintain Fitness in Rembang Regency. *Journal of Physical Education Health and Sport*, 12(1), 82–87. <https://doi.org/10.15294/JPEHS.V12I1.29016>

- Hochstetler, D. (2024). KUDOS, SEGMENTS, AND HEATMAPS: SEEKING A MEANINGFUL LIFE USING STRAVA. *Media Studies*, 15(29), 171–185. <https://doi.org/10.20901/ms.15.29.9/SUBMITTED>
- Kim, J., & Chung, J. (2024). Analysis of Service Quality in Smart Running Applications Using Big Data Text Mining Techniques. *Journal of Theoretical and Applied Electronic Commerce Research*, 19(4), 3352–3369. <https://doi.org/10.3390/JTAER19040162>
- Mulyana, A., Lestari, D., Pratiwi, D., Rohmah, N. M., Tri, N., Agustina, N. N. A., & Hefty, S. (2024). Menumbuhkan Gaya Hidup Sehat Sejak Dini Melalui Pendidikan Jasmani, Olahraga, Dan Kesehatan. *Jurnal Bintang Pendidikan Indonesia*, 2(2), 321–333. <https://doi.org/10.55606/JUBPI.V2I2.2998>
- Nam, K.-H., & Kwon, J. (2024). Analysis of the Acceptance Intention of GPS-Based Physical Exercise Applications : Focused on Strava. *Journal of Sport and Leisure Studies*, 95, 181–195. <https://doi.org/10.51979/KSSLS.2024.01.95.181>
- Rianti, A., Majid, N. W. A., & Fauzi, A. (2023). CRISP-DM: Metodologi Proyek Data Science. *Prosiding Seminar Nasional Teknologi Informasi Dan Bisnis (SENATIB) 2023*.
- Rivers, D. J. (2020). Strava as a discursive field of practice: Technological affordances and mediated cycling motivations. *Discourse, Context & Media*, 34, 100345. <https://doi.org/10.1016/j.dcm.2019.100345>
- Robinson, P., Johnson, P., & Vernooy, M. (2024). Strava Metro Data: How can urban planning leverage crowdsourced fitness activity data? *Canadian Planning and Policy*, 2024(1), 90–108. <https://doi.org/10.24908/CPP-APC.V2024I1.16889>
- Sopa, I. S., & Pomohaci, M. (2018). DEVELOPING A HEALTHY LIFESTYLE OF STUDENTS THROUGH THE PRACTICE OF SPORT ACTIVITIES. *Land Forces Academy Review*, XXIII(3(91)), 207. <https://doi.org/10.2478/raft-2018-0025>
- Sun, Y. (2017). EXPLORING POTENTIAL OF CROWDSOURCED GEOGRAPHIC INFORMATION IN STUDIES OF ACTIVE TRAVEL AND HEALTH: STRAVA DATA AND CYCLING BEHAVIOUR. *The International Archives of the Photogrammetry, Remote Sensing and Spatial Information Sciences*, XLII-2/W7, 1357–1361. <https://doi.org/10.5194/isprs-archives-XLII-2-W7-1357-2017>
- Venter, Z. S., Gundersen, V., Scott, S. L., & Barton, D. N. (2023). Bias and precision of crowdsourced recreational activity data from Strava. *Landscape and Urban Planning*, 232, 104686. <https://doi.org/10.1016/J.LANDURBPLAN.2023.104686>