

DAFTAR ISTILAH

<i>Accuracy</i>	: Persentase prediksi yang benar dari total jumlah data yang diuji.
<i>Activation Function</i>	: Fungsi matematis dalam neuron jaringan saraf untuk menentukan output berdasarkan input, misalnya sigmoid atau ReLU.
<i>Annotated Data</i>	: Data yang telah diberi label (misal <i>aware/not aware</i>) sebagai acuan untuk pelatihan atau evaluasi model.
<i>Awareness</i>	: Tingkat kesadaran atau pemahaman individu terhadap suatu isu, dalam konteks ini adalah kesehatan mental.
<i>Aware</i>	: Label untuk komentar yang menunjukkan kesadaran terhadap isu kesehatan mental.
<i>Bag of Words</i>	: Representasi teks sebagai kumpulan kata tanpa memperhatikan urutan atau konteks kata.
<i>Batch Size</i>	: Jumlah sampel data yang diproses sekaligus sebelum memperbarui parameter model.
<i>Bidirectional LSTM</i>	: LSTM yang memproses data dalam dua arah—maju dan mundur—untuk menangkap konteks yang lebih luas.
<i>Classification Report</i>	: Laporan hasil evaluasi model klasifikasi, berisi metrik seperti accuracy, precision, recall, dan F1-score.
<i>Class Imbalance</i>	: Ketidakseimbangan jumlah data antara dua atau lebih kelas dalam dataset.
<i>Comment Corpus</i>	: Sekumpulan komentar teks dari TikTok yang dianalisis dalam penelitian.
<i>Confusion Matrix</i>	: Matriks evaluasi performa model klasifikasi yang menggambarkan perbandingan antara hasil prediksi dan kondisi aktual.
<i>Corpus</i>	: Sekumpulan dokumen teks yang digunakan sebagai bahan analisis atau pelatihan model.
<i>Cross-Validation</i>	: Teknik evaluasi model dengan membagi data ke dalam beberapa subset untuk menguji generalisasi model.
<i>Dropout</i>	: Teknik regulasi dalam neural network untuk mencegah overfitting dengan menghilangkan sebagian neuron secara acak saat pelatihan.
<i>Early Stopping</i>	: Mekanisme untuk menghentikan pelatihan model secara otomatis ketika nilai parameter tertentu tidak membaik lagi.
<i>Embedding Layer</i>	: Lapisan dalam neural network untuk mengubah kata menjadi vektor berdimensi tetap.
<i>Epoch</i>	: Satu iterasi penuh ketika seluruh dataset digunakan dalam pelatihan model.
<i>Evaluation Metrics</i>	: Ukuran atau indikator performa model, seperti accuracy, recall, precision, dan F1-score.
<i>FastText</i>	: Model <i>word embedding</i> yang mempertimbangkan sub-kata (subword) untuk menangani kata-kata tidak dikenal.

<i>Feature Extraction</i>	:	Proses mengekstraksi informasi penting dari data mentah untuk digunakan dalam pelatihan model.
<i>F1-Score</i>	:	Harmonic mean dari precision dan recall; menggambarkan keseimbangan antara keduanya.
FN (<i>False Negative</i>)	:	Prediksi model yang salah mengklasifikasikan data positif sebagai negatif.
FP (<i>False Positive</i>)	:	Prediksi model yang salah mengklasifikasikan data negatif sebagai positif.
Gensim	:	<i>Library Python untuk NLP, terutama untuk topic modelling seperti LDA.</i>
<i>Hyperparameter</i>	:	Parameter eksternal model yang ditentukan sebelum pelatihan, seperti jumlah neuron, learning rate, dll.
KDD (<i>Knowledge Discovery in Database</i>)	:	Proses eksplorasi dan ekstraksi pengetahuan dari data yang besar dan kompleks.
<i>Labelling</i>	:	Proses memberi anotasi atau label pada data, misalnya komentar TikTok diberi label aware atau not aware.
LDA (<i>Latent Dirichlet Allocation</i>)	:	Algoritma probabilistik untuk mengidentifikasi topik dalam kumpulan dokumen teks.
<i>Learning Rate</i>	:	Kecepatan pembaruan bobot dalam proses pelatihan model neural network.
<i>Lexicon</i>	:	Kumpulan kata atau frasa yang digunakan sebagai dasar pelabelan berdasarkan kamus kata.
<i>Loss Function</i>	:	Fungsi yang digunakan untuk mengukur kesalahan antara prediksi model dan label aktual.
LSTM (<i>Long Short-Term Memory</i>)	:	Jenis jaringan saraf berulang (RNN) yang mampu menyimpan informasi dalam jangka waktu panjang, cocok untuk data teks berurutan.
<i>Not Aware</i>	:	Label untuk komentar yang tidak menunjukkan kesadaran terhadap isu kesehatan mental.
<i>Overfitting</i>	:	Kondisi ketika model terlalu cocok dengan data pelatihan hingga gagal melakukan generalisasi terhadap data baru.
<i>Preprocessing</i>	:	Proses pembersihan dan transformasi data mentah sebelum dianalisis lebih lanjut.
<i>Precision</i>	:	Rasio antara jumlah prediksi positif yang benar dibandingkan dengan total prediksi positif.
PyLDAVis	:	Tool visualisasi untuk model LDA yang menampilkan <i>intertopic distance map</i> .
<i>Recall</i>	:	Rasio antara jumlah prediksi positif yang benar dibandingkan dengan total data positif aktual.
<i>Regularization</i>	:	Teknik untuk mencegah overfitting dengan menambahkan penalti dalam fungsi loss.
RNN (<i>Recurrent Neural Network</i>)	:	Jaringan saraf yang digunakan untuk memproses data sekuensial dengan mempertimbangkan urutan input.
<i>Sigmoid</i>	:	Fungsi aktivasi yang digunakan pada output layer untuk menghasilkan probabilitas antara 0 dan 1.

<i>Stopwords</i>	: Kata-kata umum yang dihilangkan dalam pemrosesan teks karena dianggap tidak memiliki makna penting dalam konteks analisis (misal: "dan", "yang").
<i>Text Normalization</i>	: Proses merapikan teks dari bentuk tidak baku menjadi bentuk standar, seperti menghapus simbol atau angka.
<i>Tokenization</i>	: Proses memecah teks menjadi unit terkecil seperti kata atau token.
<i>Topic Modelling</i>	: Teknik untuk mengidentifikasi struktur tematik tersembunyi dari kumpulan dokumen teks.
TP (<i>True Positive</i>)	: Data positif yang berhasil diklasifikasikan dengan benar sebagai positif oleh model.
TN (<i>True Negative</i>)	: Data negatif yang berhasil diklasifikasikan dengan benar sebagai negatif oleh model.
<i>Tuning</i>	: Proses penyesuaian hyperparameter untuk mendapatkan performa model terbaik.
<i>Val Accuracy</i>	: Akurasi yang diperoleh saat model diuji terhadap data validasi, digunakan untuk memantau overfitting.
<i>Wordcloud</i>	: Visualisasi kata dalam bentuk awan di mana ukuran kata menunjukkan frekuensi kemunculannya.