

BAB V

PENUTUP

5.1 Kesimpulan

Berdasarkan hasil analisis data dan pembahasan yang telah dilakukan terhadap sampel penelitian sebanyak 100 responden pengguna media sosial usia 17–30 tahun, dengan data uji berupa sepuluh pasangan konten stimulus foto (10 foto asli dan 10 foto hasil rekayasa *Deepfake AI*), mengenai hubungan antara Pengetahuan tentang *Deepfake AI* (X1) dengan Kemampuan Deteksi (Y1), dapat ditarik kesimpulan sebagai berikut:

1. Tingkat Kemampuan Masyarakat dalam Membedakan Konten Nyata dan *Deepfake AI*

Tingkat kemampuan masyarakat (responden) dalam membedakan antara konten asli dan rekayasa *Deepfake AI* berada pada kategori cukup tinggi, dengan tingkat akurasi rata-rata keseluruhan sebesar 77,1% dari 10 butir soal uji. Distribusi skor variabel Kemampuan Deteksi (Y1) memuncak pada skor 8 dari total skor maksimum 10 (dengan frekuensi terbanyak sekitar 35 responden), serta nilai estimasi rata-rata (constant) regresi sebesar 8,051. Temuan ini menandakan bahwa secara visual dan intuitif, masyarakat memiliki ketajaman persepsi dasar dalam mengenali hal yang tidak biasa atau kejanggalan pada media digital.

2. Hasil Uji Korelasi antara Pengetahuan dan Kemampuan Deteksi

Pengujian Korelasi Pearson Product Moment menunjukkan nilai koefisien korelasi $r = -0,028$ dengan nilai signifikansi $p\text{-value} = 0,784$ ($p > 0,05$). Hasil ini membuktikan bahwa tidak terdapat hubungan yang signifikan antara Pengetahuan tentang *Deepfake AI* (X_1) dengan Kemampuan Deteksi (Y_1), dengan tingkat kekuatan hubungan yang tergolong sangat lemah.

3. Hasil Uji Regresi dan Pengujian Hipotesis

Model regresi $\hat{Y} = 8,051 - 0,017X_1$ terbukti tidak layak (not fit) dengan nilai $F = 0,075$ ($p = 0,784$) dan nilai koefisien determinasi (R^2) sebesar 0,001, yang berarti Pengetahuan tentang *Deepfake AI* (X_1) hanya memberikan kontribusi sebesar 0,1% terhadap Kemampuan Deteksi (Y_1), sedangkan 99,9% sisanya dipengaruhi oleh faktor lain di luar model. Berdasarkan nilai signifikansi $p = 0,784 > 0,05$, disimpulkan bahwa H_0 DITERIMA dan H_a DITOLAK, artinya pengetahuan teoritis mengenai *Deepfake AI* yang dimiliki seseorang tidak secara otomatis meningkatkan keterampilan praktisnya dalam mendeteksi konten rekayasa *Deepfake*.

4. Kelebihan Metode/Pendekatan yang Digunakan

Pendekatan deskriptif kuantitatif dengan kuesioner daring (*Google Forms*) yang digunakan dalam penelitian ini memiliki beberapa kelebihan. Dari segi objektivitas, hasil diperoleh dari perhitungan statistik sehingga meminimalkan bias subjektif peneliti. Dari segi efisiensi, pendekatan ini mampu menjangkau 100 responden dalam waktu relatif singkat tanpa perlu mendatangi lokasi responden secara langsung. Dari segi kemudahan

pengolahan dan komunikasi hasil, data dapat diolah secara cepat menggunakan bahasa pemrograman *Python* dan disajikan dalam bentuk tabel maupun grafik yang mudah dipahami, seperti pada Tabel 4.2, Tabel 4.4, dan Gambar 4.10. Selain itu, pendekatan ini juga dapat direplikasi oleh peneliti lain menggunakan prosedur pengambilan data yang sama, sehingga memperkuat tingkat kepercayaan terhadap hasil.

5. Kelemahan Metode/Pendekatan yang Digunakan

Di sisi lain, pendekatan ini juga memiliki beberapa kelemahan. Kurangnya kedalaman eksplorasi menjadi salah satu kelemahan, karena data yang dihasilkan berupa angka sehingga kurang mampu menggali alasan atau proses berpikir responden secara mendalam saat menilai konten *Deepfake*. Selain itu, terdapat potensi bias responden dan rendahnya kontrol atas kondisi pengisian kuesioner daring, seperti kesungguhan menjawab, gangguan lingkungan, maupun keterbatasan pemahaman responden terhadap pertanyaan. Keterbatasan stimulus penelitian juga menjadi kelemahan, karena hanya menggunakan 10 sampel foto sehingga belum mampu merepresentasikan seluruh variasi kualitas dan jenis konten *Deepfake* (video maupun audio) yang ada di dunia nyata. Kelemahan terakhir adalah sifat data yang dikumpulkan pada satu titik waktu (*cross-sectional*), sehingga tidak dapat menggambarkan perubahan kemampuan deteksi responden dari waktu ke waktu.

5.2 Saran

Berdasarkan temuan dan kesimpulan penelitian ini, termasuk kelemahan pendekatan yang telah diuraikan pada Sub Bab 5.1 poin 5, diajukan beberapa saran sebagai berikut:

1. Solusi untuk Mengatasi Kelemahan Metode/Pendekatan Penelitian

Untuk mengatasi keterbatasan kedalaman eksplorasi, penelitian selanjutnya disarankan mengombinasikan kuesioner kuantitatif dengan wawancara atau *focus group discussion* (pendekatan mixed methods) agar dapat menggali alasan dan proses berpikir responden secara lebih mendalam. Untuk mengurangi potensi bias responden pada kuesioner daring, disarankan menambahkan pertanyaan kontrol (*attention check*) serta memperjelas instruksi pengisian. Untuk mengatasi keterbatasan jumlah dan jenis stimulus, disarankan memperbanyak dan memvariasikan sampel uji (foto, audio, dan video) dengan tingkat kecanggihan *Deepfake* yang beragam. Terakhir, untuk mengatasi sifat data *cross-sectional*, disarankan menggunakan desain penelitian longitudinal agar perkembangan kemampuan deteksi masyarakat dapat diamati dari waktu ke waktu.

2. Bagi Masyarakat

Masyarakat diimbau untuk tidak hanya berpuas diri pada pemahaman konseptual/teoritis mengenai *Deepfake AI*, melainkan melatih diri dengan keterampilan verifikasi praktis. Masyarakat perlu membiasakan tindakan verifikasi silang (*fact-checking*), memvalidasi sumber berita resmi, menggunakan alat penelusuran gambar terbalik (*reverse image search*), serta

bersikap kritis dan skeptis terhadap media visual yang berpotensi memicu sensasi emosional.

3. Bagi Pemerintah dan Pembuat Kebijakan

Pemerintah melalui kementerian dan lembaga terkait (seperti Kominfo) disarankan untuk mereformulasi strategi dan kurikulum literasi digital. Kampanye literasi digital hendaknya tidak lagi berfokus pada sosialisasi teori semata, melainkan dialihkan ke pelatihan simulasi praktis (*hands-on training*) berbasis kasus nyata. Selain itu, pemerintah perlu memperkuat regulasi tata kelola kecerdasan buatan serta menyediakan infrastruktur publik berupa perangkat lunak pendeteksi *Deepfake* gratis yang mudah diakses oleh masyarakat umum.

4. Bagi Peneliti Selanjutnya

Peneliti berikutnya disarankan untuk:

Memperluas cakupan wilayah geografis serta keberagaman demografi responden guna meningkatkan generalisasi hasil penelitian; menambah variasi stimulus penelitian dengan mengombinasikan berbagai bentuk konten rekayasa (*Deepfake* audio, video, serta gambar statis) dengan variasi tingkat kecanggihan algoritma; dan mempertimbangkan pendekatan *Machine Learning*, misalnya membandingkan batas kemampuan deteksi intuitif manusia (*human perception*) dengan kinerja model deteksi berbasis *Convolutional Neural Network* (CNN), untuk memperoleh gambaran komparatif yang lebih mendalam mengenai efektivitas deteksi *Deepfake*.