



Client-Side Real-Time Inference untuk Efisiensi Bandwidth: Komparasi Lightweight Super-Resolution ESPCN dan FSRCNN Berbasis Browser pada Web Retail

Muhammad Nur*, Carudin, Faizal Kurnia Ramdhana

Fakultas Teknologi Informasi dan Digital, Program Studi Teknik Informatika, Universitas Bani Saleh, Bekasi, Indonesia

Email: ^{1,*}mnur@ubs.ac.id, ²carudin@ubs.ac.id, ³faizalkurnia076@ubs.ac.id

Email Penulis Korespondensi: mnur@ubs.ac.id

Abstrak—Industri e-commerce Indonesia yang tumbuh pesat menuntut penyajian citra produk berkualitas tinggi, namun distribusi langsung citra beresolusi tinggi menimbulkan beban bandwidth yang signifikan khususnya pada infrastruktur jaringan seluler yang fluktuatif. Penelitian-penelitian sebelumnya tentang pemrosesan visual sisi klien difokuskan pada video streaming, sementara pengembangan arsitektur Lightweight Super-Resolution (LSR) dilakukan secara eksklusif dalam lingkungan komputasi native menyisakan celah empiris pada implementasi LSR di dalam peramban web untuk konteks platform web retail. Penelitian ini mengisi celah tersebut dengan mengimplementasikan dan mengomparasi dua arsitektur LSR ESPCN dan FSRCNN yang dieksekusi secara langsung di dalam peramban menggunakan TensorFlow.js dengan akselerasi WebGL pada platform web retail CV Citra Wyrus Sakti. Evaluasi dilakukan terhadap 28 sampel citra produk industri mencakup tiga dimensi: efisiensi bandwidth, waktu inferensi lintas perangkat, dan kualitas restorasi visual. Hasilnya membuktikan bahwa pendekatan ini mereduksi konsumsi bandwidth rata-rata sebesar 50,1%, yang berdampak langsung pada peningkatan skor Lighthouse Performance dari 79 menjadi 99 dan penurunan Largest Contentful Paint (LCP) dari 3,8 detik menjadi 0,8 detik. ESPCN terbukti 1,84× lebih cepat dari FSRCNN pada laptop dan 1,52× lebih cepat pada perangkat seluler, sekaligus unggul dalam merestorasi ketajaman tepi dan detail tekstur tanpa efek penghalusan berlebih. Kontribusi utama penelitian ini bersifat dua: menghadirkan salah satu perbandingan empiris pertama mengenai inferensi LSR sisi klien pada citra produk e-commerce dunia nyata, serta menyediakan kerangka implementasi berbasis peramban yang dapat digunakan kembali memadukan TensorFlow.js, akselerasi WebGL, dan alur praproses HR–LR yang dapat diadopsi platform ritel Indonesia tanpa perubahan infrastruktur di sisi native. ESPCN direkomendasikan sebagai arsitektur optimal yang menyeimbangkan efisiensi bandwidth dengan estetika visual e-commerce dalam lingkungan peramban web.

Kata Kunci: Lightweight Super-Resolution; Efisiensi Bandwidth; Client-Side Inference; ESPCN; FSRCNN; Web Performance

Abstract—Indonesia's rapidly expanding e-commerce industry demands high-quality product imagery, yet directly distributing high-resolution images imposes a significant bandwidth burden particularly on fluctuating mobile network infrastructures. Prior research on client-side visual processing has focused on video streaming, while Lightweight Super-Resolution (LSR) architecture development has been conducted exclusively in native computing environments leaving an empirical gap in browser-based LSR deployment for retail web platforms. This study bridges that gap by implementing and comparing two LSR architectures ESPCN and FSRCNN executed directly within the browser using TensorFlow.js with WebGL acceleration on the CV Citra Wyrus Sakti retail web platform. Evaluation was conducted on 28 industrial product image samples across three dimensions: bandwidth efficiency, cross-device inference time, and visual restoration quality. Results demonstrate that this approach reduces average bandwidth consumption by 50.1%, directly translating to a Lighthouse Performance score increase from 79 to 99 and a Largest Contentful Paint (LCP) reduction from 3.8 seconds to 0.8 seconds. ESPCN proved 1.84× faster than FSRCNN on laptops and 1.52× faster on mobile devices, while also outperforming in restoring edge sharpness and texture detail without excessive smoothing artifacts. The main contribution of this study is twofold: it provides one of the first empirical comparisons of client-side LSR inference on real-world e-commerce product imagery, and it delivers a reusable browser-based deployment framework combining TensorFlow.js, WebGL acceleration, and an HR–LR preprocessing pipeline that Indonesian retail platforms can adopt without native-side infrastructure changes. ESPCN is recommended as the optimal architecture for balancing bandwidth efficiency with e-commerce visual aesthetics within browser environments.

Keywords: Lightweight Super-Resolution; Bandwidth Efficiency; Client-Side Inference; ESPCN; FSRCNN; Web Performance

1. PENDAHULUAN

Indonesia merupakan salah satu pasar perdagangan elektronik terbesar di Asia Tenggara dengan GMV sebesar 52,93 miliar dolar pada 2023 dan proyeksi 86,81 miliar dolar pada 2028 [1]. Ekspansi pasar ini ditopang oleh 185,3 juta pengguna internet aktif yang mayoritas mengakses e-commerce melalui perangkat seluler [2]. Kondisi ini mendorong persaingan antarplatform ritel daring yang semakin ketat dan menjadikan kualitas pengalaman visual sebagai faktor diferensiasi yang tidak dapat diabaikan.

Dalam ekosistem perdagangan daring, konsumen tidak dapat berinteraksi secara fisik dengan produk sehingga citra produk menjadi substitusi utama yang membantu pembeli mengevaluasi tekstur, warna, bentuk, dan kualitas barang sebelum mengambil keputusan pembelian. Penelitian oleh Mokobombang dan Kusumawati [3] membuktikan bahwa foto produk berpengaruh positif dan signifikan terhadap niat beli konsumen pada platform e-commerce. Sementara itu, Yang et al. [4] mengungkapkan bahwa isyarat visual-taktil pada platform e-commerce berkontribusi pada pembentukan persepsi terhadap karakteristik produk dan secara langsung meningkatkan purchase intention. Temuan-temuan ini secara empiris menegaskan bahwa kualitas resolusi citra produk bukan sekadar aspek estetika, melainkan variabel strategis yang terukur pengaruhnya terhadap konversi dan performa bisnis platform ritel daring dan dengan demikian menjadikan kualitas pengiriman citra sebagai masalah rekayasa yang memerlukan solusi teknis yang tepat.

Paradoks yang muncul adalah semakin tinggi resolusi citra produk yang disajikan, semakin besar pula ukuran berkas yang harus ditransmisikan. Berdasarkan laporan HTTP Archive [5], aset gambar menyumbang lebih dari 60%



total bobot halaman web secara global. Bobot halaman yang berlebih secara langsung berkorelasi dengan peningkatan latensi muat, yang diukur melalui metrik Largest Contentful Paint (LCP) sebagaimana ditetapkan dalam standar Core Web Vitals [6]. Nilai LCP yang buruk tidak hanya merusak pengalaman pengguna, tetapi juga berdampak negatif pada peringkat SEO dan tingkat konversi halaman.

Permasalahan ini semakin berat di Indonesia mengingat kondisi infrastruktur internet seluler yang belum merata. Kecepatan unduh jaringan seluler Indonesia menunjukkan disparitas signifikan antarwilayah, berkisar antara 15 Mbps di wilayah timur hingga 45 Mbps di Pulau Jawa, dengan rata-rata nasional sekitar 18 Mbps berdasarkan data Speedtest by Ookla [7], [8]. Pada jaringan dengan kecepatan rendah atau fluktuatif, pengiriman citra beresolusi tinggi berukuran lebih dari 1 MB per gambar dapat menyebabkan waktu tunggu yang signifikan, mendorong pengguna untuk meninggalkan halaman sebelum konten visual termuat sepenuhnya. Pendekatan konvensional bertumpu pada kompresi citra menggunakan format seperti JPEG, PNG, atau WebP. Meskipun mampu mereduksi ukuran berkas, metode kompresi lossy secara inheren menghilangkan sebagian informasi piksel secara permanen, menghasilkan artefak blocking, penurunan ketajaman tepi, dan hilangnya detail tekstur halus [9]. Dalam konteks produk industri seperti komponen elektronik dan sensor suhu, degradasi semacam ini dapat menimbulkan ambiguitas bagi calon pembeli dalam menilai kualitas material dan finishing produk. Teknologi Super-Resolution (SR) berbasis deep learning menawarkan paradigma alternatif yang fundamental: alih-alih mengurangi informasi, SR merekonstruksi citra beresolusi tinggi dari masukan beresolusi rendah melalui pemetaan non-linear yang dipelajari secara data-driven [10]. Varian Lightweight Super-Resolution (LSR) dirancang khusus untuk meminimalkan kebutuhan komputasi sehingga proses rekonstruksi dapat dijalankan secara efisien pada perangkat dengan sumber daya terbatas [11][12].

Kemajuan ekosistem JavaScript modern, khususnya TensorFlow.js [13], telah membuka jalan bagi eksekusi model deep learning secara langsung di dalam peramban web. Dengan memanfaatkan akselerasi GPU melalui WebGL, TensorFlow.js memungkinkan operasi tensor intensif seperti konvolusi dan upsampling dieksekusi secara efisien [13]. Kajian terkini oleh Wang et al. [14] mengkonfirmasi bahwa latensi inferensi deep learning berbasis WebGL telah mencapai titik yang memungkinkan aplikasi visual nyata dalam lingkungan peramban.

Kajian terhadap penelitian terdahulu mengungkap dikotomi yang menyisakan celah penelitian yang jelas. Pertama, Goh et al. [13] mengulas pengembangan dan deployment aplikasi deep learning berbasis peramban secara umum, tetapi tidak menyentuh implementasi maupun evaluasi model super-resolution sama sekali. Kedua, Wang et al. [14] membuktikan bahwa latensi inferensi deep learning berbasis WebGL telah layak untuk aplikasi visual nyata, namun pengujian tersebut dilakukan pada model generik dan tidak spesifik pada arsitektur super-resolution ataupun kasus penggunaan e-commerce. Ketiga, Yang et al. [15] menunjukkan bahwa client-side video processing mampu meningkatkan Quality of Experience (QoE) sebesar 5–14% melalui adaptive streaming, tetapi penerapannya terbatas pada konten video bergerak dan tidak mengevaluasi citra produk statis. Keempat, penelitian arsitektur SR itu sendiri AsConvSR [11], BSRN [12], model berbasis Swin Transformer [16][17], hingga SR real-time terkuantisasi pada NPU mobile [18], seluruhnya divalidasi pada benchmark akademik generik dan dieksekusi secara native atau pada chip khusus (NPU), sehingga kelayakannya saat dijalankan di dalam runtime JavaScript peramban yang jauh lebih terbatas belum pernah diuji. Bahkan model yang lebih kompleks seperti Real-ESRGAN [19] terbukti unggul untuk degradasi dunia nyata namun terlalu berat secara komputasi untuk eksekusi sisi klien. Dengan demikian, celah penelitian yang teridentifikasi bersifat ganda: (1) belum ada studi yang mengevaluasi LSR secara spesifik di dalam lingkungan peramban berbasis WebGL, dan (2) belum ada studi yang menguji efektivitas LSR pada dataset citra produk e-commerce dunia nyata alih-alih benchmark akademik generik. Kesenjangan ganda inilah yang membedakan penelitian ini dari kelima penelitian terdahulu tersebut dan menjadi motivasi utamanya: menghadirkan bukti empiris tentang efektivitas LSR dalam peramban web pada konteks platform e-commerce retail dengan dataset gambar produk dunia nyata.

Penelitian ini bertujuan mengimplementasikan dan mengevaluasi pendekatan LSR berbasis sisi klien pada platform web retail, dengan komparasi ESPCN [16] dan FSRCNN [17] mencakup tiga dimensi: (1) besaran reduksi bandwidth, (2) kelayakan waktu inferensi pada laptop dan smartphone, serta (3) kualitas restorasi visual citra produk. Temuan penelitian ini diharapkan memberikan landasan empiris bagi pengembang platform e-commerce Indonesia dalam merancang strategi optimasi pengiriman citra yang mempertimbangkan realitas infrastruktur jaringan lokal.

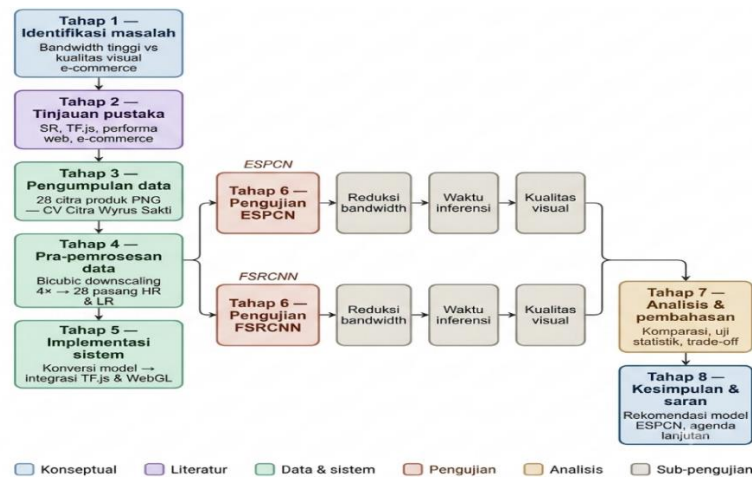
Kontribusi utama penelitian ini terhadap bidang ini mencakup tiga hal: (1) menyediakan bukti empiris mengenai kelayakan dan performa komparatif dua arsitektur LSR (ESPCN dan FSRCNN) yang dieksekusi murni di sisi klien melalui TensorFlow.js dengan akselerasi WebGL; (2) merumuskan kerangka evaluasi tiga-dimensi efisiensi bandwidth, waktu inferensi lintas perangkat, dan kualitas restorasi visual yang dapat direplikasi pada platform e-commerce lain; dan (3) menghasilkan rekomendasi arsitektur berbasis data sebagai alternatif kompresi lossy konvensional, yang mempertimbangkan realitas disparitas infrastruktur jaringan seluler di Indonesia.

2. METODOLOGI PENELITIAN

2.1 Tahapan Penelitian

Secara keseluruhan, penelitian ini dilaksanakan melalui delapan tahapan yang terstruktur dan saling berkesinambungan, sebagaimana divisualisasikan pada Gambar 1. Tahap 1 dimulai dengan identifikasi masalah yang berpusat pada paradoks antara kebutuhan kualitas visual tinggi dan keterbatasan bandwidth pada platform e-commerce Indonesia. Tahap 2 melakukan tinjauan pustaka sistematis untuk memetakan celah penelitian pada persimpangan LSR dan peramban web.

Tahap 3 mengumpulkan 28 citra produk nyata dari CV Citra Wyrus Sakti sebagai dataset. Tahap 4 membentuk pasangan citra HR-LR melalui bicubic downscaling dengan faktor 4×. Tahap 5 mengimplementasikan sistem LSR sisi klien dengan mengintegrasikan model ESPCN dan FSRCNN ke dalam TensorFlow.js. Tahap 6 menjalankan pengujian paralel pada kedua model terhadap tiga dimensi evaluasi: reduksi bandwidth, waktu inferensi, dan kualitas visual. Tahap 7 menganalisis dan membandingkan hasil kedua model secara statistik. Tahap 8 merumuskan kesimpulan dan rekomendasi berdasarkan temuan empiris yang diperoleh.



Gambar 1. Alur Tahapan Penelitian Lightweight Super-Resolution Berbasis Sisi Klien

2.2 Dataset dan Preprocessing

Dataset penelitian bersumber dari 28 citra produk milik CV Citra Wyrus Sakti, perusahaan manufaktur komponen pemanas industri dan sensor suhu di Bekasi, Jawa Barat. Citra produk dikategorikan: 10 citra aksesoris, 6 citra pemanas (heaters), dan 12 citra sensor suhu, seluruhnya dalam format PNG beresolusi 1200×1200 piksel dengan ukuran rata-rata 1.194 KB. Komposisi lengkap dataset disajikan pada Tabel 1.

Tabel 1. Komposisi Dataset Citra Produk CV Citra Wyrus Sakti

Kategori Produk	Jumlah Citra	Resolusi (px)	Ukuran Rata-rata (KB)	Format
Aksesoris	10	1200 × 1200	1.087	PNG
Pemanas (Heaters)	6	1200 × 1200	1.198	PNG
Sensor Suhu	12	1200 × 1200	1.254	PNG
Total / Rata-rata	28	1200 × 1200	1.194	PNG

Dari 28 citra HR tersebut, dibentuk pasangan citra LR menggunakan teknik Bicubic Downscaling dengan faktor skala 4× menggunakan skrip Python 3.13.5 dengan pustaka Pillow, sehingga citra beresolusi 1200×1200 piksel menghasilkan pasangan LR berukuran 300×300 piksel. Pemilihan faktor skala 4× mengikuti konvensi standar dalam evaluasi super-resolution [18] dan memberikan tantangan rekonstruksi yang signifikan. Seluruh proses menghasilkan 56 citra pengujian (28 pasang HR-LR), total ukuran dataset 31,8 MB untuk citra HR dan sekitar 16 MB untuk citra LR.

2.3 Arsitektur Sistem dan Model LSR

Sistem yang dirancang menerapkan arsitektur dua lapis. Pada lapis peladen, file HTML, CSS, JavaScript, berkas model LSR, serta citra LR disimpan dan didistribusikan. Peladen tidak melakukan pemrosesan gambar apapun. Pada lapis klien, alur pemrosesan berjalan: (1) peramban memuat halaman web dan mengunduh berkas model LSR secara asinkron; (2) model dimuat ke memori GPU melalui fungsi `tf.loadGraphModel()` dari TensorFlow.js [13]; (3) setiap citra LR dikonversi menjadi tensor input menggunakan `tf.browser.fromPixels()`; (4) inferensi dijalankan menggunakan `tf.tidy()` untuk manajemen memori yang efisien; (5) tensor output dikonversi kembali menjadi gambar menggunakan `tf.browser.toPixels()`; serta (6) seluruh proses dibungkus dalam fungsi asinkron `async/await` [19] untuk menjaga responsivitas antarmuka. ESPCN (Efficient Sub-Pixel Convolutional Neural Network) [16] beroperasi sepenuhnya dalam ruang resolusi rendah. Arsitekturnya terdiri dari beberapa lapisan konvolusi (filter 5×5 dan 3×3), diikuti satu lapisan sub-pixel convolution di bagian akhir yang melakukan upscaling melalui operasi pixel shuffle. Keunggulan kritis ESPCN adalah seluruh operasi perkalian matriks dilakukan pada tensor berukuran kecil (dimensi LR 300×300 piksel). FSRCNN (Fast Super-Resolution Convolutional Neural Network) [17] menggunakan empat tahap utama: feature extraction, shrinking, mapping, dan expanding, diikuti lapisan dekonvolusi untuk upsampling. Lapisan dekonvolusi FSRCNN beroperasi pada ruang output berukuran 1200×1200 piksel 16 kali lebih besar dalam jumlah elemen sehingga membutuhkan alokasi buffer GPU yang jauh lebih besar.

Pemilihan ESPCN dan FSRCNN didasarkan pada tiga pertimbangan utama. Pertama, keduanya merupakan arsitektur foundational LSR berbasis pure convolution yang kompatibel penuh dengan backend WebGL TensorFlow.js,



tidak seperti arsitektur modern berbasis attention atau transformer [20] yang membutuhkan operasi tensor non-standar yang belum didukung secara optimal oleh WebGL. Kedua, model-model mutakhir seperti BSRN [12], AsConvSR [11], dan Swin Transformer-based SR [20] dikembangkan dan dievaluasi dalam lingkungan native PyTorch, sehingga feasibility-nya dalam lingkungan browser belum diverifikasi. Ketiga, kesederhanaan arsitektur keduanya memungkinkan investigasi murni terhadap dampak perbedaan strategi upsampling (sub-pixel convolution vs. dekonvolusi) pada performa sisi klien. Lebih lanjut, model berbasis distilasi pengetahuan (knowledge distillation) dan varian FSRCNN yang dioptimasi secara modern belum memiliki konverter TensorFlow.js yang stabil dan terverifikasi pada saat penelitian dilaksanakan, sehingga tidak dapat dijamin kompatibilitasnya dengan backend WebGL tanpa modifikasi arsitektur yang di luar cakupan penelitian ini. Perlu juga dicatat bahwa penelitian ini tidak mengeksplorasi optimasi model melalui quantization (misalnya INT8) yang berpotensi meningkatkan kecepatan inferensi secara signifikan pada perangkat seluler kelas bawah [21], aspek ini merupakan salah satu agenda penelitian lanjutan yang direkomendasikan.

Kedua model dikonversi dari format Keras (.h5) ke format TensorFlow.js Graph Model menggunakan alat tensorflowjs_converter dengan presisi floating-point 32-bit. Spesifikasi perangkat keras pengujian disajikan pada Tabel 2 dan spesifikasi perangkat lunak disajikan pada Tabel 3.

Tabel 2. Spesifikasi Perangkat Keras Pengujian

Komponen	Laptop	Mobile
Prosesor	Intel Core i5-8265U @ 1.60–1.80 GHz	ARM Octa-core (Mid-Range)
Memori (RAM)	12 GB DDR4	8 GB + 3 GB Extended
GPU / Akselerator	NVIDIA GeForce MX230 (2 GB VRAM)	Integrated Mobile GPU
Peramban	Google Chrome v148.0.7778.167	Google Chrome Mobile

Tabel 3. Spesifikasi Perangkat Lunak Pengujian

Perangkat Lunak	Versi	Fungsi
Python	3.13.5	Preprocessing citra, pembentukan pasangan HR-LR
TensorFlow.js	4.22.0	Eksekusi inferensi model LSR di peramban
tensorflowjs_converter	4.22.0	Konversi model Keras (.h5) ke format TF.js
Google Chrome	148.0.7778.167	Peramban pengujian utama
Chrome DevTools	Bawaan Chrome	Pengukuran transfer data dan profil jaringan
Google Lighthouse	Bawaan Chrome	Audit performa website (LCP, FCP, skor performa)
Microsoft Excel	2021	Analisis statistik paired t-test

2.4 Prosedur Pengujian

Pengujian reduksi bandwidth menggunakan tab Network Chrome DevTools dengan kondisi Disable Cache aktif dan No Throttling. Persentase reduksi dihitung: $\text{Reduksi (\%)} = \frac{(\text{Ukuran_HR} - \text{Ukuran_LR})}{\text{Ukuran_HR}} \times 100\%$. Signifikansi statistik diuji menggunakan uji paired t-test dengan $\alpha = 0,05$. Pengujian performa website menggunakan Google Lighthouse [22] untuk mengaudit metrik LCP, FCP, TBT, dan CLS. Waktu inferensi diukur menggunakan fungsi performance.now() pada JavaScript [19] dengan presisi sub-milidetik. Setiap model dieksekusi sebanyak 15 iterasi (P1–P15) pada seluruh 28 sampel di masing-masing perangkat, menghasilkan 420 titik data per kombinasi model-perangkat [14]. Evaluasi kualitas visual dilakukan secara deskriptif-kualitatif terhadap lima aspek: ketajaman tepi, fidelitas tekstur, akurasi warna, kemunculan artefak, dan keterbacaan detail kecil [23].

3. HASIL DAN PEMBAHASAN

3.1 Implementasi Sistem

Implementasi sistem LSR berbasis sisi klien pada platform web CV Citra Wyrus Sakti berhasil diselesaikan dengan mengintegrasikan model ESPCN dan FSRCNN yang telah dikonversi ke format TensorFlow.js ke dalam antarmuka halaman daftar produk. Sistem dibangun menggunakan HTML5, CSS3, dan JavaScript murni tanpa framework front-end tambahan untuk meminimalkan overhead muat aset dan memaksimalkan kontrol atas alur eksekusi inferensi. Antarmuka pengujian menyediakan empat tombol kontrol mode tampilan: Mode HR (citra PNG beresolusi penuh sebagai baseline), Mode LR (citra LR tanpa rekonstruksi), Mode ESPCN, dan Mode FSRCNN. Pergantian mode dapat dilakukan secara instan tanpa reload halaman, memungkinkan perbandingan visual secara langsung. Model dimuat saat halaman pertama kali diakses menggunakan mekanisme preloading asinkron sehingga latensi cold start tidak dirasakan pengguna pada saat interaksi aktif.

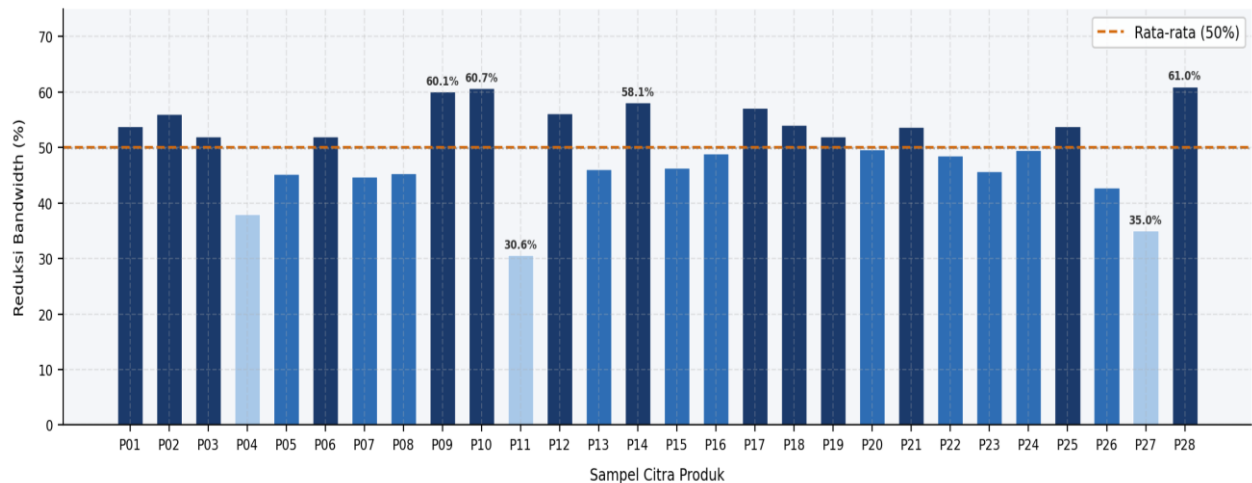
3.2 Hasil Pengukuran Reduksi Bandwidth

Pengukuran terhadap seluruh 28 citra produk pada kedua skenario menghasilkan data yang menunjukkan perbedaan ukuran transmisi yang konsisten dan signifikan. Tabel 4 menyajikan data ukuran file untuk sampel-sampel yang dipilih secara representatif beserta persentase reduksi yang dicapai.

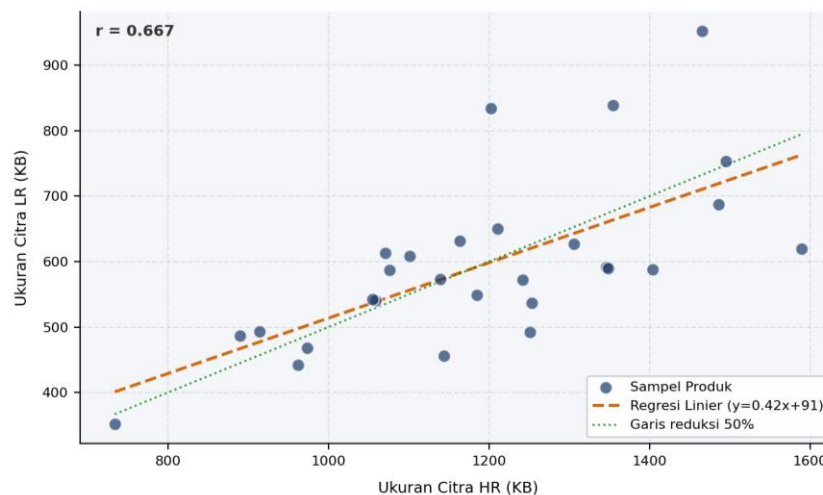
Tabel 4. Ukuran Transmisi Citra HR dan LR pada Sampel Representatif

No	Nama File	HR (KB)	LR (KB)	Reduksi (%)	Kategori
P01	product01.png	1.486	687	53,8	Aksesori
P06	product06.png	734	352	52,0	Pemanas
P10	product10.png	1.251	492	60,7	Sensor Suhu
P11	product11.png	1.202	834	30,6	Aksesori
P27	product27.png	1.465	952	35,0	Aksesori
P28	product28.png	1.589	620	61,0	Sensor Suhu
Rata-rata	—	1.194	596	50,1	Semua kategori
Total (28 produk)	—	33.420	16.680	50,1	—

Data pada Tabel 4 memperlihatkan karakteristik distribusi yang kaya. Rentang ukuran HR membentang dari 734 KB (P06) hingga 1.589 KB (P28). Setelah downscaling bicubic $4\times$, ukuran berkas yang ditransmisikan turun drastis. Reduksi tidak seragam, sampel P10 dan P28 mencapai 61%, sementara P11 dan P27 hanya 30–35%. Pola ini mencerminkan hubungan antara kompleksitas frekuensi tinggi pada citra dan tingkat kompresibilitas PNG: citra dengan area warna homogen menghasilkan kompresi lebih efisien, sebaliknya citra bertekstur rapat menghasilkan file LR relatif lebih besar.

**Gambar 2.** Distribusi Persentase Reduksi Bandwidth per Sampel Citra Produk (n=28)

Gambar 2 memvisualisasikan distribusi persentase reduksi bandwidth pada seluruh 28 sampel. Garis putus-putus horizontal menandai nilai rata-rata 50%. Distribusi memperlihatkan dua kelompok: mayoritas sampel (sekitar 20 dari 28) berkumpul di rentang 45–62%, sementara sekelompok kecil berada di rentang 30–44%, yang mengindikasikan variasi alami akibat perbedaan kompleksitas tekstur antarkategori produk. Gambar 3 menyajikan plot sebar korelasi antara ukuran HR dan LR. Koefisien korelasi Pearson ($r = 0,667$) mengindikasikan hubungan linear positif yang sedang. Garis regresi linier ($y = 0,46x + 47$) menunjukkan bahwa setiap penambahan 1 KB pada ukuran HR rata-rata hanya menambah sekitar 0,46 KB pada ukuran LR, mengkonfirmasi penghematan relatif yang konsisten [9].

**Gambar 3.** Korelasi Ukuran File HR dan LR pada 28 Sampel Produk



3.3 Validasi Statistik Paired T-Test

Signifikansi perbedaan ukuran bandwidth antara skenario HR dan LR diuji menggunakan paired t-test terhadap 28 pasangan data. Hipotesis yang diuji: $H_0: \mu_{HR} = \mu_{LR}$ (tidak ada perbedaan) versus $H_1: \mu_{HR} \neq \mu_{LR}$ (terdapat perbedaan signifikan). Hasil pengujian statistik secara lengkap disajikan pada Tabel 5.

Tabel 5. Hasil Uji Paired T-Test Reduksi Bandwidth HR dan LR (n=28)

Parameter Uji	Nilai	Interpretasi
Mean HR (KB)	1.193,57	Rata-rata ukuran citra beresolusi penuh
Mean LR (KB)	595,57	Rata-rata ukuran citra beresolusi rendah
Selisih Mean (KB)	598,00	Penghematan rata-rata per citra
Standar Deviasi Selisih	151,32	Variabilitas penghematan antarsampel
Jumlah Pasangan (n)	28	Jumlah citra produk yang diuji
t-statistik	20,92	Nilai uji t yang dihitung
Derajat Kebebasan (df)	27	$n - 1$
p-value (two-tail)	$3,25 \times 10^{-18}$	Jauh di bawah $\alpha = 0,05$
Tingkat Signifikansi (α)	0,05	Ambang kepercayaan 95%
Keputusan	H_0 ditolak	Perbedaan signifikan secara statistik

Hasil pada Tabel 5 menunjukkan nilai t-statistik sebesar 20,92 dengan p-value sebesar $3,25 \times 10^{-18}$ jauh di bawah ambang signifikansi $\alpha = 0,05$. Hipotesis nol ditolak dengan keyakinan yang sangat tinggi. Standar deviasi selisih sebesar 151,32 KB mengindikasikan bahwa meskipun besaran penghematan bervariasi antarsampel, arah penghematan konsisten secara menyeluruh. Implikasi operasional dari penghematan 598 KB per citra sangat signifikan: platform dengan 10.000 citra yang diakses 1.000 kali per hari akan menghemat sekitar 598 GB data transfer harian.

3.4 Analisis Performa Website

Pengujian performa website menggunakan Google Lighthouse memberikan perspektif yang lebih operasional tentang dampak pendekatan LSR. Perbandingan metrik performa secara komprehensif disajikan pada Tabel 6.

Tabel 6. Perbandingan Metrik Performa Website Mode HR dan Mode LR

Metrik Performa	Mode HR	Mode LR	Perubahan
Total Transfer Data	~37 MB	~16 MB	-56,8%
Skor Lighthouse Performance	79 / 100	99 / 100	+20 poin
First Contentful Paint (FCP)	0,2 detik	0,3 detik	+0,1 dtk
Largest Contentful Paint (LCP)	3,8 detik	0,8 detik	-78,9%
Total Blocking Time (TBT)	90 ms	80 ms	-11,1%
Cumulative Layout Shift (CLS)	0	0	Tidak berubah
Kategori LCP	Needs Improvement	Good	Naik kategori

Data pada Tabel 6 mengungkap kontras yang tajam. Peningkatan skor Lighthouse dari 79 menjadi 99 merepresentasikan lompatan dari kategori "Needs Improvement" ke "Good" perubahan yang berdampak langsung pada peringkat halaman dalam hasil pencarian Google, mengingat Core Web Vitals menjadi faktor dalam algoritma ranking mesin pencari [22]. Penurunan LCP dari 3,8 detik menjadi 0,8 detik (reduksi 78,9%) adalah temuan paling impresif. Nilai LCP 3,8 detik pada mode HR mendekati ambang "Poor" (di atas 4,0 detik), sementara 0,8 detik pada mode LR berada jauh di dalam zona "Good" (di bawah 2,5 detik).

3.5 Analisis Waktu Inferensi

Waktu inferensi diukur melalui 15 iterasi berulang pada 28 sampel di dua perangkat, menghasilkan 420 titik data per kombinasi model-perangkat [14]. Tabel 7 merangkum statistik deskriptif komprehensif dari keseluruhan pengukuran.

Tabel 7. Statistik Deskriptif Waktu Inferensi ESPCN dan FSRCNN per Perangkat (n=420 per sel)

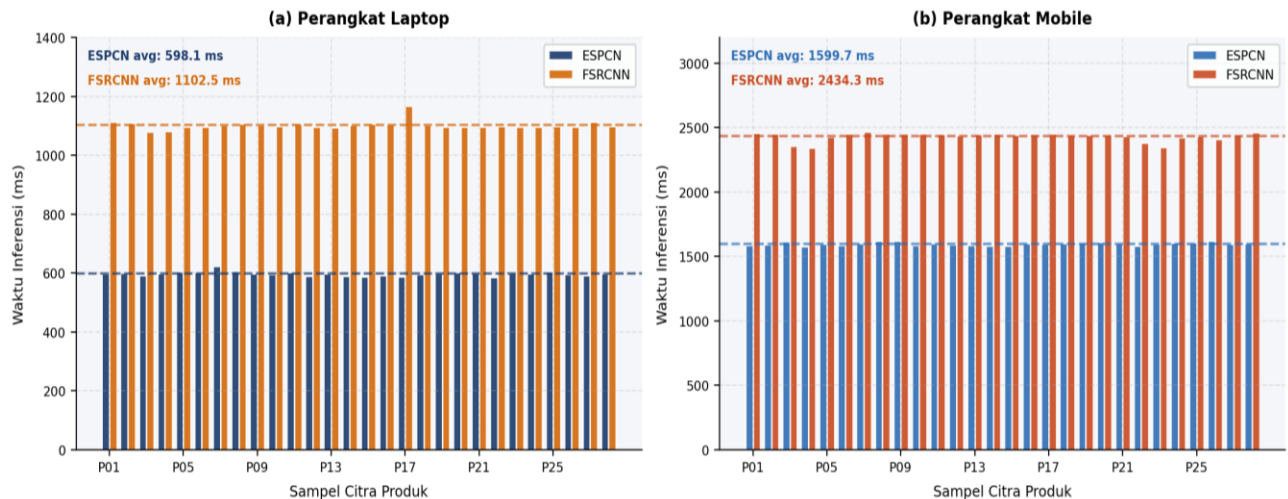
Parameter	ESPCN Laptop	FSRCNN Laptop	ESPCN Mobile	FSRCNN Mobile
Rata-rata (ms)	598,14	1.102,51	1.599,70	2.434,28
Minimum (ms)	574,40	1.060,00	1.506,40	2.313,60
Maksimum (ms)	717,50	2.119,70	1.683,50	2.521,00
Standar Deviasi (ms)	6,80	48,63	5,08	11,42
Koefisien Variasi (%)	1,14	4,41	0,32	0,47
CI 95% ($\pm$ ms)	2,92	5,79	4,74	13,20
Rasio vs ESPCN	1,00×	1,84×	—	1,52×

Tabel 7 dan Gambar 4 mengungkap dua pola utama. Pertama, ESPCN secara konsisten lebih cepat dari FSRCNN di kedua perangkat: 1,84× lebih cepat pada laptop (598,14 ms vs 1.102,51 ms) dan 1,52× pada smartphone (1.599,70 ms

vs 2.434,28 ms). Kedua, ESPCN menunjukkan stabilitas jauh lebih baik: koefisien variasi 1,14% vs 4,41% untuk FSRCNN pada laptop, mengindikasikan prediktabilitas latensi yang tinggi.

Perbedaan performa ini berakar pada arsitektur masing-masing model saat dikompilasi melalui WebGL backend TensorFlow.js [13]. ESPCN [16] melakukan seluruh operasi konvolusi pada tensor 300×300 piksel (dimensi LR), dengan upsampling melalui `tf.depthToSpace` yang merupakan operasi memory-shuffle ringan. FSRCNN [17] menggunakan lapisan dekonvolusi yang pada saat eksekusi di WebGL diterjemahkan menjadi operasi konvolusi transposa pada tensor output 1200×1200 piksel 16 kali lebih besar dalam jumlah elemen.

Analisis bottleneck arsitektur FSRCNN dapat dipahami lebih dalam melalui perspektif alokasi memori GPU. Pada backend WebGL TensorFlow.js, setiap tensor operasional dialokasikan sebagai tekstur GPU. ESPCN memproses seluruh konvolusi pada tensor 300×300 piksel, menghasilkan alokasi buffer yang relatif kecil per lapisan. Sebaliknya, lapisan dekonvolusi FSRCNN mengoperasikan tensor output berukuran 1200×1200 piksel mengharuskan alokasi buffer GPU yang jauh lebih besar selama eksekusi. Tekanan memori GPU yang lebih tinggi ini tidak hanya memengaruhi waktu inferensi rata-rata, tetapi juga berkontribusi pada koefisien variasi yang jauh lebih tinggi pada FSRCNN (4,41% vs 1,14% pada laptop), yang mencerminkan ketidakstabilan latensi akibat manajemen tekstur GPU yang lebih kompleks. Secara kuantitatif, tensor output 1200×1200 piksel pada FSRCNN menghasilkan alokasi buffer sekitar 16 kali lebih besar dibandingkan ESPCN (1.440.000 piksel vs 90.000 piksel pada ruang LR), yang berarti setiap lapisan dekonvolusi FSRCNN mengonsumsi memori tekstur GPU secara proporsional lebih besar sebuah overhead yang nyata terutama pada perangkat mobile dengan memori GPU terbatas dan pengelolaan konteks WebGL yang lebih ketat [13][14]. Meskipun latensi pada perangkat mobile (~1.600 ms untuk ESPCN) lebih terasa, proses berlangsung secara asinkron tanpa memblokir interaksi pengguna.



Gambar 4. Perbandingan Waktu Inferensi ESPCN dan FSRCNN pada Perangkat Laptop (a) dan Mobile (b)

3.6 Evaluasi Kualitas Visual

Evaluasi visual dilakukan dengan membandingkan empat versi citra secara berdampingan: HR (acuan), LR (input sistem), SR-ESPCN (output yang diusulkan), dan SR-FSRCNN (pembanding). Tabel 8 merangkum hasil evaluasi kualitatif pada lima dimensi.

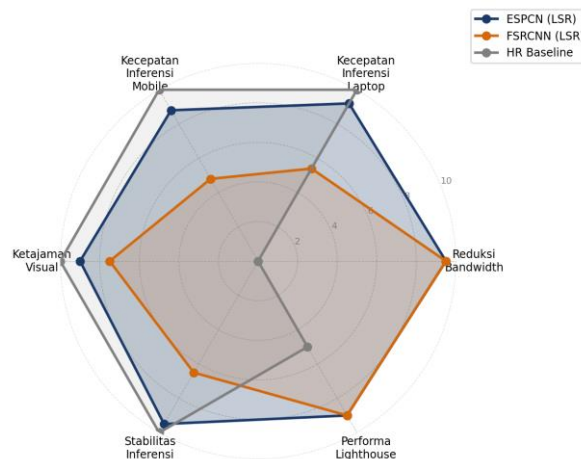
Tabel 8. Matriks Evaluasi Kualitas Visual Citra pada Empat Kondisi

Dimensi Evaluasi	Citra LR	SR-ESPCN	SR-FSRCNN	HR (Acuan)
Ketajaman tepi objek	Rendah buram	Tinggi	Sedang agak halus	Referensi
Fidelitas tekstur	Hilang homogen	Terjaga baik	Sebagian halus	Referensi
Akurasi warna	Pudar, kontras rendah	Natural, kontras tinggi	Sedikit pudar	Referensi
Artefak visual	Blur parah, pikselasi	Minimal, tidak mengganggu	Smoothing pada detail	Tidak ada
Keterbacaan detail kecil	Rendah	Tinggi	Sedang	Referensi
Kesesuaian overall dengan HR	Sangat jauh	Mendekati HR	Cukup dekat	100%

Tabel 8 memperlihatkan pola yang konsisten: ESPCN mengungguli FSRCNN pada seluruh dimensi yang dievaluasi [23]. Keunggulan paling mencolok terlihat pada dimensi ketajaman tepi dan fidelitas tekstur. Pada sampel Band & Strip Heaters (zoom-in), ESPCN berhasil mengembalikan tepi lingkaran komponen, garis batas bodi logam, dan konektor dengan ketajaman yang hampir setara dengan citra HR. Detail kisi-kisi metalik yang hilang pada versi LR

tampak kembali pada rekonstruksi ESPCN. Sebaliknya, rekonstruksi FSRCNN menghasilkan tepi yang lebih lembut akibat efek smoothing.

Pada sampel Cartridge Heater, perbedaan menjadi semakin dramatis. Rekonstruksi ESPCN berhasil memisahkan susunan kabel individual dengan garis pemisah yang jelas, sementara rekonstruksi FSRCNN menampilkan efek penggabungan (merging) antarkabel yang membuat susunan tampak lebih padat. Bagi calon pembeli yang ingin menilai kualitas insulasi dan susunan kabel pada produk heater industri, perbedaan ini memiliki dampak praktis yang nyata terhadap kemampuannya mengevaluasi produk.



Gambar 5. Radar Chart Perbandingan Performa Komprehensif ESPCN, FSRCNN, dan HR Baseline

Gambar 5 merangkum trade-off performa antar-skenario dalam visualisasi multidimensi. ESPCN mendominasi pada hampir semua dimensi dengan area poligon yang lebih luas dibandingkan FSRCNN. Satu-satunya dimensi di mana keduanya setara adalah efisiensi bandwidth (keduanya menggunakan distribusi citra LR yang sama). HR Baseline hanya unggul pada dimensi kualitas visual referensi namun tertinggal jauh pada bandwidth dan performa loading. Visualisasi ini mengkonfirmasi bahwa ESPCN menawarkan keseimbangan terbaik di antara semua opsi yang diuji [23].

3.7 Pembahasan: Trade-Off, Implikasi, dan Posisi Penelitian

Mengintegrasikan seluruh temuan, penelitian ini mengungkapkan proposisi nilai yang kuat: pengalihan beban transmisi ke komputasi sisi klien bukan sekadar substitusi teknis yang netral, melainkan sebuah peningkatan menyeluruh yang menguntungkan semua pihak. Peladen mendapat beban transmisi lebih ringan, pengguna mendapat halaman yang dimuat lebih cepat, dan kualitas visual yang diterima setelah rekonstruksi selesai jauh melampaui apa yang bisa dicapai oleh kompresi lossy konvensional [9] pada ukuran file yang setara. Perlu dicatat bahwa trade-off ini tidak sepenuhnya bebas biaya: FCP justru mengalami peningkatan kecil dari 0,2 detik menjadi 0,3 detik pada mode LR (Tabel 6), mencerminkan overhead pemuatan model LSR saat halaman pertama kali diakses. Namun trade-off ini bersifat asimetris, pengguna menerima sedikit lebih lambat pada momen pertama, sebaliknya mendapat LCP yang 78,9% lebih cepat sebagai kompensasi yang jauh melampaui biayanya.

Temuan penelitian ini secara langsung mengkonfirmasi dan memperluas klaim Wang et al. [14] yang menyatakan bahwa latensi inferensi WebGL telah mencapai titik layak untuk aplikasi visual nyata. Penelitian ini tidak hanya memverifikasi klaim tersebut, tetapi juga mengkuantifikasikannya dalam konteks spesifik: 598 ms pada laptop dan 1.600 ms pada smartphone mid-range merupakan angka konkret yang dapat digunakan pengembang sebagai referensi perencanaan UX. Dibandingkan dengan Yang et al. [15] yang melaporkan peningkatan Quality of Experience sebesar 5–14% melalui client-side processing pada video streaming, pendekatan ini menghasilkan penghematan bandwidth 50,1%, jauh melampaui skala yang dilaporkan untuk video meskipun perbandingan langsung tidak sepenuhnya setara karena perbedaan domain konten (gambar statis vs. video bergerak). Di sisi lain, arsitektur ESPCN dan FSRCNN yang diuji merupakan fondasi dari model-model mutakhir seperti BSRN [12] dan AsConvSR [11], sehingga validasi keduanya dalam lingkungan browser memberikan landasan untuk mengeksplorasi turunan arsitektur yang lebih canggih di masa depan.

Dari perspektif arsitektur sistem, temuan ini mendukung paradigma edge computing dan client-side processing untuk mengurangi beban infrastruktur terpusat [26]. Relevansi untuk konteks Indonesia secara spesifik dapat dikuantifikasi: pada kecepatan unduh rata-rata 18 Mbps di tingkat nasional [7], penghematan 598 KB per citra setara dengan pengurangan waktu transmisi sekitar 0,27 detik per gambar dan pada jaringan seluler di wilayah timur dengan kecepatan 15 Mbps [8], penghematan tersebut mencapai ~0,32 detik per gambar. Pada halaman produk yang menampilkan 10–20 gambar secara bersamaan, efek kumulatif penghematan ini bersifat signifikan dan langsung dirasakan pengguna. Kontribusi penelitian ini terletak pada pengisian celah empiris tersebut: membuktikan bahwa arsitektur LSR dapat beroperasi secara efektif dalam lingkungan peramban web menggunakan TensorFlow.js pada kondisi dataset dan perangkat dunia nyata.



Keunggulan ESPCN atas FSRCNN dalam konteks browser bukan semata soal kecepatan, melainkan mencerminkan kesesuaian arsitektural yang fundamental dengan model komputasi WebGL. Backend WebGL pada TensorFlow.js mengalokasikan setiap tensor sebagai tekstur GPU dan mengeksekusi operasi melalui shader programs sebuah model komputasi yang dirancang untuk memproses data dalam jumlah besar secara paralel pada resolusi tetap. ESPCN, yang melakukan seluruh operasinya pada dimensi LR sebelum upsampling akhir, memanfaatkan karakteristik ini secara optimal: buffer GPU yang kecil dan konsisten menghasilkan manajemen shader yang efisien. Sebaliknya, lapisan dekonvolusi FSRCNN yang beroperasi pada dimensi output penuh (1200×1200 piksel) memaksa WebGL untuk mengalokasikan dan mengelola tekstur berukuran 16 kali lebih besar, yang berdampak bukan hanya pada waktu eksekusi tetapi juga pada fragmentasi memori GPU fenomena yang menjelaskan mengapa koefisien variasi FSRCNN (4,41%) jauh lebih tinggi dari ESPCN (1,14%). Temuan ini memberikan prinsip desain praktis: untuk deployment LSR berbasis browser, arsitektur yang menjaga operasi pada ruang resolusi rendah (pre-upsampling) akan secara inheren lebih kompatibel dengan model komputasi WebGL dibandingkan arsitektur yang beroperasi di ruang resolusi tinggi.

Satu keterbatasan metodologis yang perlu diakui secara eksplisit adalah sifat evaluasi kualitas visual yang masih deskriptif-kualitatif tanpa metrik kuantitatif seperti PSNR (Peak Signal-to-Noise Ratio), SSIM (Structural Similarity Index), atau LPIPS (Learned Perceptual Image Patch Similarity). Pendekatan kualitatif dipilih dengan pertimbangan bahwa dalam konteks e-commerce, persepsi visual konsumen bukan nilai piksel absolut adalah yang menentukan keputusan pembelian [23]. Namun demikian, absennya metrik kuantitatif membatasi kemampuan membandingkan hasil secara langsung dengan benchmark SR yang ada dalam literatur, yang sebagian besar menggunakan dataset Set5, Set14, atau BSD100 dengan metrik PSNR/SSIM. Keterbatasan kedua adalah jumlah sampel (28 citra) yang berasal dari satu domain produk industri spesifik komponen pemanas dan sensor suhu yang memiliki karakteristik tekstur (tepi tajam, detail metalik, latar belakang homogen) yang tidak merepresentasikan keseluruhan keragaman visual kategori produk e-commerce Indonesia.

Mempertimbangkan keseluruhan kontribusi dan keterbatasannya, penelitian ini secara spesifik mengisi posisi yang belum ditempati dalam lanskap riset: yakni sebagai studi empiris pertama yang memverifikasi kelayakan LSR berbasis browser menggunakan dataset gambar produk e-commerce dunia nyata bukan dataset benchmark sintesis dalam kondisi deployment yang nyata pada platform yang beroperasi. Berbeda dengan penelitian SR yang mengevaluasi model dalam lingkungan native PyTorch yang terkontrol [11][12] [20] [24], atau penelitian client-side processing yang berfokus pada infrastruktur tanpa mengevaluasi SR secara spesifik [13] [15], penelitian ini menjembatani keduanya. Posisi ini sekaligus mendefinisikan batas klaim yang dapat dibuat: penelitian ini membuktikan kelayakan teknis dan efektivitas bandwidth, tetapi belum membuktikan penerimaan pengguna akhir (yang memerlukan studi MOS), skalabilitas enterprise (yang memerlukan pengujian beban), maupun keunggulan visual terukur atas kompresi konvensional (yang memerlukan uji PSNR/SSIM). Agenda riset lanjutan yang paling kritis adalah validasi ketiga aspek tersebut secara bertahap.

4. KESIMPULAN

Penelitian ini telah membuktikan secara empiris bahwa pendekatan Lightweight Super-Resolution berbasis client-side inference, menggunakan TensorFlow.js–WebGL pada platform web retail nyata merupakan strategi yang layak secara teknis untuk mengatasi paradoks antara kualitas visual dan efisiensi bandwidth pada e-commerce Indonesia. Dari tiga dimensi evaluasi, ESPCN terbukti sebagai arsitektur optimal: unggul dalam efisiensi komputasi, stabilitas latensi, maupun kualitas restorasi visual dibandingkan FSRCNN, yang bottleneck-nya berakar pada beban alokasi memori GPU dari lapisan dekonvolusinya. Kontribusi utama penelitian ini terhadap literatur adalah menghadirkan salah satu bukti empiris pertama mengenai kelayakan dan performa komparatif LSR sisi klien pada kasus nyata e-commerce ritel Indonesia, sekaligus kerangka evaluasi tiga-dimensi yang dapat diadopsi oleh pengembang platform retail lain. Signifikansi praktis temuan ini melampaui satu platform: pengembang e-commerce Indonesia, yang beroperasi dalam realitas disparitas jaringan seluler antarwilayah, kini memiliki landasan empiris untuk mengadopsi strategi distribusi citra LR+rekonstruksi sisi klien sebagai alternatif kompresi lossy konvensional, dengan potensi penghematan bandwidth skala besar yang telah terverifikasi secara statistik. Namun demikian, beberapa keterbatasan mendasar perlu diakui agar hasil penelitian ini tidak diinterpretasikan melampaui cakupannya: (1) evaluasi kualitas visual bersifat deskriptif tanpa metrik kuantitatif seperti PSNR, SSIM, atau LPIPS; (2) dataset terbatas pada 28 citra dari satu domain produk industri spesifik, sehingga generalisabilitas ke kategori produk e-commerce yang lebih luas memerlukan validasi tersendiri; (3) optimasi quantization (INT8) yang berpotensi krusial untuk perangkat seluler kelas bawah belum dieksplorasi; dan (4) yang paling kritis, skalabilitas sistem untuk mengelola ribuan citra secara konkuren pada heterogenitas ekstrem spesifikasi perangkat pengguna Indonesia yang sesungguhnya (dari flagship hingga entry-level) belum diuji, sehingga performa sistem di luar kondisi pengujian terkontrol ini masih belum dapat dijamin. Penelitian lanjutan direkomendasikan untuk mengeksplorasi arsitektur hybrid CNN-Transformer yang kompatibel WebGL, penambahan metrik PSNR/SSIM/LPIPS, validasi pengguna melalui Mean Opinion Score (MOS), serta integrasi CDN untuk deployment skala enterprise.

REFERENCES

- [1] P. Putra Dharma, M. S. Purwanegara, and S. A. Wibowo, "Proposed Marketing Strategy to Increase Brand Loyalty: Study Case of Lazada Indonesia," *International Journal of Current Science Research and Review*, vol. 7, no. 7, pp. 4701–4716, 2024, doi: 10.47191/ijcsrr/v7-i7-15.



- [2] W. A. Social and Meltwater, "Digital 2024: Indonesia," DataReportal, 2024. [Online]. Available: https://wearesocial.com/id/wp-content/uploads/sites/19/2024/02/Digital_2024_Indonesia.pdf
- [3] M. G. Mokobombang and N. Kusumawati, "Impact of Product Description, Product Photo, Rating, and Review on Purchase Intention in E-Commerce," *Journal of Consumer Studies and Applied Marketing*, vol. 1, no. 2, pp. 137–147, 2023, doi: 10.58229/jcsam.v1i2.100.
- [4] Y. Yang, Q. Yang, and X. Liu, "Seeing the Feel, Willing to Buy: How Visual–Tactile Cues Shape Consumer Purchase Intention in E-Commerce Platforms," *Journal of Theoretical and Applied Electronic Commerce Research*, vol. 21, no. 3, p. 84, Mar. 2026, doi: 10.3390/jtaer.21030084.
- [5] B. Dornauer and M. Felderer, "Web Image Formats: Assessment of Their Real-World-Usage and Performance Across Popular Web Browsers," in *Product-Focused Software Process Improvement*, vol. 14483, R. Kadgien, A. Jedlitschka, A. Janes, V. Lenarduzzi, and X. Li, Eds., Cham: Springer, 2024, pp. 132–147. doi: 10.1007/978-3-031-49266-2_9.
- [6] K. Król and W. Sroka, "Internet in the Middle of Nowhere: Performance of Geoportals in Rural Areas According to Core Web Vitals," *ISPRS Int. J. Geoinf.*, vol. 12, no. 12, p. 484, Nov. 2023, doi: 10.3390/ijgi12120484.
- [7] Patrick Welay, Caroline Herfisia Parlina, Ivena Julia Armelinda, and Jadianan Parhusip, "The Distribution Of Average Mobile Network Speeds In Indonesia Based On Ookla Data," *Informatika: Jurnal Teknik Informatika dan Multimedia*, vol. 4, no. 2, pp. 24–29, Dec. 2024, doi: 10.51903/informatika.v4i2.825.
- [8] N. Salsi Anugrah, H. Adilah, D. P. Panggabean, and R. P. Laksana, "Analisis Kecepatan Internet Seluler Di Indonesia Berdasarkan Regional," *Journal of Data Analytics, Information, and Computer Science*, vol. 2, no. 2, pp. 155–158, Apr. 2025, doi: 10.70248/jdaics.v2i2.1821.
- [9] V. I. Ungureanu, P. Negirla, and A. Korodi, "Image-Compression Techniques: Classical and Region-of-Interest-Based Approaches Presented in Recent Papers," *Sensors*, vol. 24, no. 3, p. 791, 2024, doi: 10.3390/s24030791.
- [10] A. Liu, Y. Liu, J. Gu, Y. Qiao, and C. Dong, "Blind Image Super-Resolution: A Survey and Beyond," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 45, no. 5, pp. 5461–5480, 2023, doi: 10.1109/TPAMI.2022.3203009.
- [11] J. Guo *et al.*, "AsConvSR: Fast and Lightweight Super-Resolution Network with Assembled Convolutions," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops (CVPRW)*, 2023, pp. 1582–1592. doi: 10.1109/CVPRW59228.2023.00160.
- [12] Z. Li *et al.*, "Blueprint Separable Residual Network for Efficient Image Super-Resolution," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops (CVPRW)*, 2022, pp. 832–842. doi: 10.1109/CVPRW56347.2022.00099.
- [13] H. A. Goh, C. C. Ho, and F. S. Abas, "Front-end deep learning web apps development and deployment: a review," *Applied Intelligence*, vol. 53, no. 12, pp. 15923–15945, 2023, doi: 10.1007/s10489-022-04278-6.
- [14] Q. Wang *et al.*, "Anatomizing Deep Learning Inference in Web Browsers," *ACM Transactions on Software Engineering and Methodology*, vol. 34, no. 2, 2025, doi: 10.1145/3688843.
- [15] J. Yang, Y. Jiang, and S. Wang, "Enhancement or Super-Resolution: Learning-Based Adaptive Video Streaming with Client-Side Video Processing," in *IEEE International Conference on Communications (ICC)*, 2022, pp. 739–744. doi: 10.1109/ICC45855.2022.9839262.
- [16] B. B. Moser, F. Raue, S. Frolov, S. Palacio, J. Hees, and A. Dengel, "Hitchhiker's Guide to Super-Resolution: Introduction and Recent Advances," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 45, no. 8, pp. 9862–9882, Aug. 2023, doi: 10.1109/TPAMI.2023.3243794.
- [17] J. Li *et al.*, "A Systematic Survey of Deep Learning-Based Single-Image Super-Resolution," *ACM Comput. Surv.*, vol. 56, no. 10, pp. 1–40, Oct. 2024, doi: 10.1145/3659100.
- [18] M. Saeed Rad, T. Yu, C. Musat, H. Kemal Ekenel, B. Bozorgtabar, and J.-P. Thiran, "Benefiting from Bicubically Down-Sampled Images for Learning Real-World Image Super-Resolution," in *Proceedings of the IEEE Winter Conference on Applications of Computer Vision (WACV)*, 2021, pp. 1589–1598. doi: 10.1109/WACV48630.2021.00163.
- [19] M. Haverbeke, *Eloquent JavaScript: A Modern Introduction to Programming*, 4th ed. No Starch Press, 2022. [Online]. Available: <https://eloquentjavascript.net>
- [20] T. Jing, C. Liu, and Y. Chen, "A Lightweight Single-Image Super-Resolution Method Based on the Parallel Connection of Convolution and Swin Transformer Blocks," *Applied Sciences*, vol. 15, no. 4, p. 1806, 2025, doi: 10.3390/app15041806.
- [21] A. Ignatov, R. Timofte, M. Denna, and A. A. R. Younes, "Real-Time Quantized Image Super-Resolution on Mobile NPUs, Mobile AI 2021 Challenge: Report," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops (CVPRW)*, 2021, pp. 2525–2534. doi: 10.1109/CVPRW53098.2021.00286.
- [22] P. Walton and B. Pollard, "Optimize Largest Contentful Paint," Web Dev. Accessed: Jul. 12, 2026. [Online]. Available: <https://web.dev/articles/optimize-lcp>
- [23] W. Zhou and Z. Wang, "Quality Assessment of Image Super-Resolution: Balancing Deterministic and Statistical Fidelity," in *Proceedings of the 30th ACM International Conference on Multimedia (MM)*, 2022, pp. 934–942. doi: 10.1145/3503161.3547899.
- [24] B. Möller, L. Görnhardt, and T. Fingscheidt, "A Lightweight Image Super-Resolution Transformer Trained on Low-Resolution Images Only," *Procedia Comput. Sci.*, vol. 264, pp. 299–308, 2025, doi: 10.1016/j.procs.2025.07.141.