

Publication Date: 02 September 2026

Archs Sci. (2026) Volume 75, Issue 6 Pages 94-100, Paper ID 2025611.
<https://doi.org/10.68304/as/75611>

Selecting Singular Components for Image Reconstruction: Exact Gradient Errors and Pixel-Fidelity Trade-offs

Shin Min Kang^{1,2}, Waqas Nazeer³

¹Department of Mathematics and Research Institute of Natural Science, Gyeongsang National University, Jinju, 52828, Korea

²Center for General Education, China Medical University, Taichung, 40402, Taiwan

³Department of Mathematics, GC University, Lahore 54000, Pakistan

Corresponding authors: Shin Min Kang (e-mail: smkang@gnu.ac.kr).

Abstract Truncated singular value decomposition minimizes squared pixel error, but selecting a different subset of the same singular components can lower spatial-gradient error. For valid forward differences, we prove that both residual errors are sums of contributions from the omitted components. Consequently, a normalized pixel–gradient objective is minimized by retaining the highest-scoring components under a fixed total budget. A two-component example shows that gradient-only selection can discard the constant intensity term. We compare four rules at five budgets on twelve photographic crops (240 reconstructions), with a separate 84-setting weight sweep and 600 exhaustive small-matrix checks. At an equivalent rank of 32 per channel, equal weighting reduces mean normalized gradient error by 0.685% relative to pixel-optimal selection, while mean normalized pixel error increases by 4.155% and mean paired PSNR decreases by 0.208 dB. On the coffee crop, gradient-only selection loses 19.723 dB of PSNR for a 1.768% gradient-error reduction. The sorting guarantee holds for the fixed, unclipped singular dictionary; the measurements show that small gradient gains may carry appreciable pixel costs.

Index Terms singular value decomposition, component selection, discrete image gradients, rank allocation, image reconstruction, reproducible numerical analysis

I. Introduction

Low-rank image reconstruction normally keeps the singular components with the largest singular values. This choice minimizes squared pixel error at a fixed rank [1], but a component with modest pixel energy can carry substantial spatial variation. If the singular vectors have already been computed and only a fixed number of dyads may be retained, the question is whether selecting components for their gradient contribution improves a reconstruction enough to justify its pixel cost.

Derivative-sensitive approximation is related to singular decompositions in Sobolev spaces [2], [3]. Here the setting is narrower: the singular vectors and values remain unchanged, and selection is made from this fixed dictionary. Randomized low-rank methods address how to obtain an approximate subspace [4]; singular-value thresholds address estimation under noise assumptions [5]. Neither determines the best subset of already-computed dyads for the residual measured here.

Our objective is squared error in signed, valid forward differences, rather than a perceptual image-quality measure. It differs from gradient-magnitude similarity [6]; we also report pixel error, PSNR, and structural similarity (SSIM) [7]. No single measure establishes perceived quality [8]. Forward differences can be viewed as incidence operators, making the residual gradient energy a quadratic form [9]. In each differen-

tiated singular dyad, the untouched factor remains orthogonal across components, yielding an exact additive identity.

We derive the resulting component score and its pixel–gradient trade-off, give an example in which gradient-only selection loses a constant intensity term, and compare four policies on twelve specified photographs. The experiment measures raw and clipped reconstructions separately so that the mathematical selection guarantee is not confused with the appearance of a displayed image.

II. Materials and Methods

A. Image Collection and Preprocessing

The reconstruction experiment uses twelve photographs distributed with scikit-image [10]: five RGB images and seven grayscale images, or 22 channel matrices in total. They include objects, faces, repeated textures, smooth surfaces, and outdoor scenes. The collection is purposive, not a random sample. Table 1 lists every input, and no image was removed on the basis of a method’s performance. The motorcycle photograph comes from the Middlebury stereo collection [11]; we use only the bundled left view, without its disparity data.

Each input is a central 256×256 crop. Its zero-based origin is obtained by flooring half the difference between the native dimension and 256. We perform no resizing or upsampling. Eight-bit intensities are divided by 255 and stored in double

Table 1: Photographic inputs and deterministic crop origins

Image identifier	Native height	Native width	Channels	Crop origin (y, x)
Astronaut	512	512	3	(128, 128)
Brick	512	512	1	(128, 128)
Camera	512	512	1	(128, 128)
Chelsea	300	451	3	(22, 97)
Clock	300	400	1	(22, 72)
Coffee	400	600	3	(72, 172)
Coins	303	384	1	(23, 64)
Grass	512	512	1	(128, 128)
Gravel	512	512	1	(128, 128)
Moon	512	512	1	(128, 128)
Motorcycle	500	741	3	(122, 242)
Rocket	427	640	3	(85, 192)

precision; RGB channels are processed separately, without luminance conversion, gamma linearization, or mean subtraction. The manifest records native dimensions, crop origins, filenames, and SHA-256 hashes for the distributed files and saved crops. These hashes identify the exact arrays used in the experiment.

B. Fixed Singular Dictionary and Residual Objectives

Let $X = (X_1, \dots, X_C)$ contain C image channels, each in $\mathbb{R}^{H \times W}$. A reduced singular value decomposition supplies

$$\begin{aligned} X_c &= \sum_{i=1}^n \sigma_{ci} u_{ci} v_{ci}^T, \\ n &= \min(H, W), \\ \sigma_{c1} &\geq \dots \geq \sigma_{cn} \geq 0. \end{aligned} \quad (1)$$

A retained set S contains channel–position pairs (c, i) . The reconstruction $\hat{X}(S)$ uses those dyads with their singular values and vectors unchanged, and $R(S) = X - \hat{X}(S)$ is its residual. We write the total budget as $B = Cr$, where r is the equivalent rank per channel. Global selection can distribute these B components unequally across channels. Thus r specifies a common representation budget, rather than the actual rank of every channel or an achieved compression ratio.

Let D_H be the $(H - 1) \times H$ matrix of valid forward differences and define D_W analogously. The two residual energies are

$$\begin{aligned} P(R) &= \sum_{c=1}^C \|R_c\|_F^2, \\ G(R) &= \sum_{c=1}^C \left(\|D_H R_c\|_F^2 + \|R_c D_W^T\|_F^2 \right). \end{aligned} \quad (2)$$

The differences include only neighboring pixels inside the crop, with no padding or scaling by physical pixel spacing. Accordingly, P measures squared intensity error, while G measures squared error in signed horizontal and vertical differences. The residual argument matters: a large value of $G(X)$ indicates substantial spatial variation, whereas a small value of $G(R)$ indicates accurate reconstruction of that variation.

For each singular component, define its intensity and gradient contributions as

$$p_{ci} = \sigma_{ci}^2, \quad d_{ci} = \sigma_{ci}^2 \left(\|D_H u_{ci}\|_2^2 + \|D_W v_{ci}\|_2^2 \right). \quad (3)$$

Unlike p_{ci} , the gradient contribution depends on the variation of the singular vectors as well as the singular value. Since a forward-difference operator has squared operator norm at most four, $0 \leq d_{ci} \leq 8p_{ci}$. This bound does not impose the same ordering on the two contributions. A smaller, more oscillatory component can have a larger gradient contribution than a smooth component with greater pixel energy.

Proposition 1 (Exact residual accounting). *For every retained subset of a fixed singular dictionary,*

$$P(R(S)) = \sum_{(c,i) \notin S} p_{ci}, \quad G(R(S)) = \sum_{(c,i) \notin S} d_{ci}. \quad (4)$$

Proof. The intensity identity follows from Frobenius orthogonality of the singular dyads. For the vertical derivative, the inner product between two differentiated dyads contains

$$\langle (D_H u_{ci}) v_{ci}^T, (D_H u_{cj}) v_{cj}^T \rangle_F = \langle D_H u_{ci}, D_H u_{cj} \rangle \langle v_{ci}, v_{cj} \rangle. \quad (5)$$

When $i \neq j$, the unchanged right singular vectors are orthogonal, so this cross term vanishes. In the horizontal derivative, orthogonality of the unchanged left singular vectors gives the corresponding cancellation. Summing the remaining diagonal terms over omitted indices and channels gives the second identity. $\square$

Applying the argument to retained components also gives $G(\hat{X}(S)) = \sum_{(c,i) \in S} d_{ci}$. Retained and omitted gradient energies therefore sum to $G(X)$ within this representation family. The unchanged orthogonal factor is essential to the proof. Arbitrary spatial weighting or nonlinear processing can remove that cancellation, so the identity must be re-established before using such modifications.

C. Selection under a Common Component Budget

For an image with $P(X) > 0$ and $G(X) > 0$, choose a weight $0 \leq \lambda \leq 1$ and minimize

$$J_\lambda(S) = (1 - \lambda) \frac{P(R(S))}{P(X)} + \lambda \frac{G(R(S))}{G(X)}, \quad |S| = B. \quad (6)$$

The denominators normalize the errors by the input's pixel and gradient energies. This makes the objective invariant to a common nonzero intensity scaling, but does not make its two terms perceptually equivalent. Both denominators include all channels, allowing the optimizer to move components between them. For a nonzero constant image, the implementation falls back to pixel-energy ordering; an all-zero image uses deterministic index ordering. All twelve crops have positive values of both energies.

Define the component score

$$q_{ci}(\lambda) = (1 - \lambda) \frac{p_{ci}}{P(X)} + \lambda \frac{d_{ci}}{G(X)}. \quad (7)$$

Proposition 2 (Optimal fixed-dictionary selection). *A set containing the B largest values of $q_{ci}(\lambda)$ minimizes J_λ over all B -component subsets of the chosen singular dictionary.*

Proof. Eq. (4) gives $J_\lambda(S) = 1 - \sum_{(c,i) \in S} q_{ci}(\lambda)$. If a retained score is smaller than an omitted score, exchanging the two increases the retained sum and decreases the objective. Repeating such exchanges gives a set of the B largest scores. Equal scores may be exchanged without changing the objective, so deterministic tie handling preserves optimality. $\square$

We compare four policies. Equal-channel prefix truncation keeps the first r components of each channel. Pixel-optimal selection keeps the B largest values of p_{ci} across all channels. Equal-weight selection uses $\lambda = 0.5$, and gradient-only selection uses $\lambda = 1$. Comparing the first two policies isolates channel allocation; comparing the latter two with pixel-optimal selection isolates the effect of the objective. The coefficient 0.5 is a specified comparison setting, rather than a weight fitted to these images.

The guarantee is conditional on the computed singular dictionary. Paired sign changes leave each dyad and its score unchanged, but a rotation within a repeated-singular-value subspace can alter both the gradient contributions and the selected dyads. Exact score ties are resolved by stable descending sorting, followed by channel index and singular position. This convention makes selection deterministic for a given dictionary; it does not resolve the basis ambiguity of a repeated singular value.

Corollary 1 (Ordered distortion trade-off). *For optimal selections at $0 \leq \lambda_1 < \lambda_2 \leq 1$, the normalized intensity residual cannot decrease and the normalized gradient residual cannot increase as the weight changes from λ_1 to λ_2 .*

Proof. Let $p_k = P(R(S_k))/P(X)$ and $g_k = G(R(S_k))/G(X)$ for an optimum S_k at λ_k . With $\Delta p = p_2 - p_1$ and $\Delta g = g_2 - g_1$, optimality gives

$$\Delta p + \lambda_1(\Delta g - \Delta p) \geq 0, \quad \Delta p + \lambda_2(\Delta g - \Delta p) \leq 0. \quad (8)$$

Subtracting yields $\Delta g - \Delta p \leq 0$. The first inequality then gives $\Delta p \geq 0$, and the second gives $\Delta g \leq (1 - \lambda_2)(\Delta g - \Delta p) \leq 0$. The argument includes both endpoints and allows equality. $\square$

The corollary concerns the two raw quadratic errors, not SSIM or visual preference. Because each score is affine in λ , component identities change only at score crossings. The selected reconstruction is constant between crossings, subject to the tie convention. Lines connecting sampled weights in a figure therefore indicate the order of evaluation, not interpolated reconstructions.

D. A Two-Component Limit on Derivative-Only Selection

An exact example clarifies why a favorable gradient objective need not preserve an image's intensity level. For $a > b > 0$, consider

$$X = \frac{a}{2} \begin{pmatrix} 1 & 1 \\ 1 & 1 \end{pmatrix} + \frac{b}{2} \begin{pmatrix} 1 & -1 \\ -1 & 1 \end{pmatrix}. \quad (9)$$

The two displayed terms are orthogonal singular dyads with squared intensity contributions a^2 and b^2 . Their gradient contributions are zero and $4b^2$, respectively. With one retained component, pixel-optimal selection keeps the constant term and incurs gradient error $4b^2$. Gradient-only selection keeps the alternating term and incurs zero gradient error, while its intensity error rises from b^2 to a^2 . Dividing the entire image by $a+b$ places its entries inside the unit interval without changing either normalized comparison. Thus the failure is not caused by an inadmissible reference intensity range.

The scores cross at $\lambda = (a^2 - b^2)/(2a^2) < 1/2$. Consequently, equal weighting can also discard the constant component: including a pixel term does not by itself protect image brightness. This example establishes a limitation of the objective without requiring a photographic test case. The experiments below examine how strongly the same problem appears in the selected crops.

E. Numerical Protocol and Implementation

The primary experiment combines twelve crops, five equivalent ranks (8, 16, 32, 64, 128), and four policies, giving 240 reconstructions. At rank 32, a separate sweep evaluates $\lambda \in \{0, 0.1, 0.25, 0.5, 0.75, 0.9, 1\}$ for all images, giving 84 settings. Some coincide with primary settings. The unit of comparison remains the parent image, not the number of reconstructions. No parameters are learned, and we use descriptive paired comparisons without significance tests or population-level confidence intervals.

Array operations use NumPy [12]; reduced singular decompositions use SciPy's LAPACK `gesdd` driver [13]. The recorded analysis environment is Python 3.13.5, NumPy 2.3.5, SciPy 1.17.0, scikit-image 0.26.0, Pillow 12.3.0, and Matplotlib 3.10.8. Linear algebra is restricted to one thread. Each channel is decomposed once, and all policies reuse its complete factors.

For each selection, we explicitly reconstruct the image and calculate its residual energies independently of the component sums. Algebraic discrepancies are divided by the larger of one and the predicted energy; this is a scaled discrepancy, rather than a relative error when the energy is below one. We also use ten standard-normal matrices at each of four shapes (3×5 , 5×3 , 4×4 , and 6×7). At five weights, every nonempty, nonfull budget is checked by exhaustive subset enumeration. The random seed is 20250917, and the resulting 600 checks test the selection implementation and the energy identities.

Peak signal-to-noise ratio (PSNR) is calculated from the unclipped reconstruction with a unit intensity range and mean squared error over all pixels and channels. SSIM is evaluated

after clipping to $[0, 1]$, using Gaussian weights, standard deviation 1.5, an eleven-pixel window, and population covariance normalization. The routine averages RGB scores across channels. SSIM uses floating-point clipped arrays; only the displayed photographs are rounded to eight-bit values. We also record the fraction of channel values outside $[0, 1]$ before clipping and recalculate gradient error after clipping.

We report arithmetic means for normalized pixel error, normalized gradient error, and SSIM. PSNR is summarized by its median because near-exact reconstruction of a low-rank crop can produce an unusually large decibel value. Where a mean PSNR loss is reported, it is the mean of paired image-level differences, not the difference of two medians. All image-level measurements are supplied as comma-separated files.

Computing the component scores from the singular vectors costs $O(Cn(H + W))$, followed by sorting Cn scores. The dense singular decomposition of a square channel remains cubic in its dimension. Storage of B unquantized dyads requires approximately $B(H + W + 1)$ scalars, so a reduced component count need not yield a smaller file than the original pixel array. Numerical figures are generated with Matplotlib [14]. The photographic panels show the clipped reconstructions at the measured 256×256 resolution.

III. Results

A. Component Contributions and Numerical Identities

In the coffee image, the pixel and gradient contributions have different distributions across singular positions (Figure 1). Panel (a) sums the three channel contributions at each position and normalizes each curve by its own total. Panel (b) shows the positions retained by each policy at $B = 96$. Equal-weight and gradient-only selection can retain later components while omitting earlier ones; their selections cannot be described by a single cutoff in each channel.

The retained-position map distinguishes channel allocation from selection within a channel. Different row totals indicate a change in allocation; gaps within a row indicate omitted singular positions. These are separate decisions. The energy curves are also normalized separately, so their vertical separation compares the distributions of the two energies, not their absolute units.

Direct reconstruction agrees with the component formulas to double-precision accuracy. Across the 240 primary reconstructions, the largest scaled discrepancies are 1.146×10^{-15} for pixel error and 1.751×10^{-15} for gradient error; the retained gradient-energy discrepancy is at most 8.289×10^{-16} . All 600 exhaustive checks return the same recorded minimum as score selection. These checks support the implementation of the identities and the selection rule.

B. Accuracy Across the Five Budgets

Both residuals decrease as the budget increases (Figure 2; Table 2). The gradient curves are close, but the gradient-only policy has a larger pixel cost at small budgets. At $r = 8$, its mean normalized pixel error is 0.381134, compared with 0.029169 for pixel-optimal selection. Their median PSNR

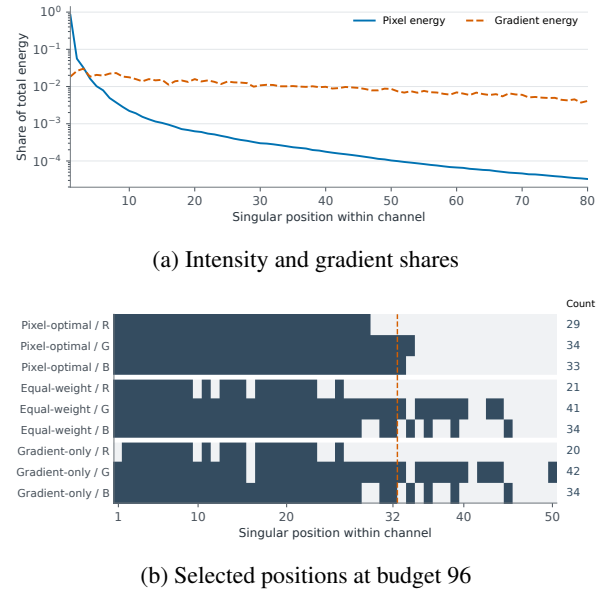


Figure 1: Pixel and gradient contributions in the coffee image. (a) Energy shares are summed across RGB channels at each singular position and normalized separately; the vertical axis is logarithmic. (b) Filled cells identify retained components at $B = 96$ ($r = 32$); the dashed line marks the equal-channel prefix boundary. Each row reports its retained count

values are 18.91 and 21.38 dB, respectively. The contrast between the mean pixel fraction and median PSNR reflects the uneven severity of the errors across images.

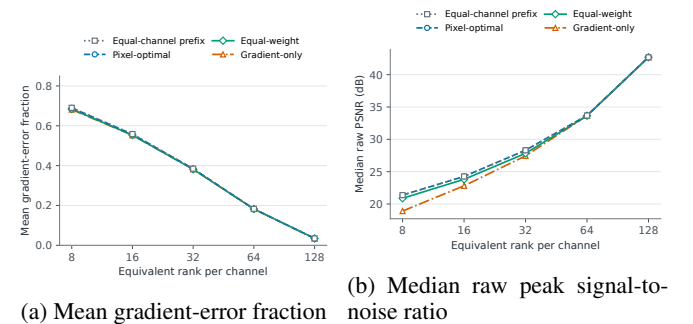


Figure 2: Budget dependence across all twelve crops. (a) Arithmetic mean of $G(R)/G(X)$. (b) Median raw PSNR. Each point summarizes the same twelve inputs, with total budget $B = Cr$. Line styles and symbols distinguish policies where curves overlap; connecting lines do not represent additional evaluated budgets

At $r = 32$, equal-weight selection reduces mean gradient error from 0.384297 to 0.381664, relative to pixel-optimal selection. The reduction is 0.685% of the pixel-optimal mean, or approximately 0.263 percentage points. Mean pixel error increases from 0.007912 to 0.008240 (4.155%), and the mean paired PSNR loss is 0.208 dB. The two policies therefore exchange a modest gradient gain for a measurable pixel penalty.

Table 2: Reconstruction accuracy across twelve crops. Pixel and gradient fractions use raw reconstructions; SSIM uses clipped floating-point images. PSNR is the image-level median, while the other columns are arithmetic means. The budget is $B = Cr$

r	Selection	Mean $P(R)/P(X)$	Mean $G(R)/G(X)$	Median PSNR (dB)	Mean SSIM
8	Equal-channel prefix	0.029221	0.690072	21.38	0.6465
8	Pixel-optimal	0.029169	0.689788	21.38	0.6466
8	Equal-weight	0.030321	0.684175	20.88	0.6452
8	Gradient-only	0.381134	0.681244	18.91	0.4689
16	Equal-channel prefix	0.016974	0.558464	24.26	0.7347
16	Pixel-optimal	0.016961	0.557238	24.26	0.7346
16	Equal-weight	0.017894	0.553176	23.81	0.7282
16	Gradient-only	0.216792	0.552059	22.81	0.6304
32	Equal-channel prefix	0.007926	0.384587	28.30	0.8357
32	Pixel-optimal	0.007912	0.384297	28.32	0.8359
32	Equal-weight	0.008240	0.381664	27.83	0.8316
32	Gradient-only	0.053087	0.381550	27.44	0.8116
64	Equal-channel prefix	0.002561	0.182881	33.68	0.9299
64	Pixel-optimal	0.002558	0.182594	33.70	0.9299
64	Equal-weight	0.002568	0.182163	33.65	0.9293
64	Gradient-only	0.002571	0.182162	33.65	0.9291
128	Equal-channel prefix	0.000370	0.034648	42.68	0.9860
128	Pixel-optimal	0.000370	0.034568	42.71	0.9860
128	Equal-weight	0.000370	0.034531	42.70	0.9860
128	Gradient-only	0.000370	0.034531	42.70	0.9860

Channel reallocation has a smaller effect. At $r = 32$, the mean normalized pixel errors of equal-channel prefix and pixel-optimal selection differ by approximately 0.000014. These policies coincide for the seven grayscale images, so the difference comes entirely from the five color images. Policy differences also narrow at larger budgets: at $r = 128$, mean gradient-error fractions range from 0.034531 to 0.034648.

C. Image-Specific Trade-Offs and Photographic Consequences

The image-level results vary across the twelve photographs (Table 3). At $r = 32$, equal weighting reduces gradient error by zero for brick and by 1.698% for clock. The paired raw PSNR decrease ranges from zero to 0.971 dB. This spread is relevant when interpreting the mean improvement.

Table 3: Paired effects of equal-weight selection relative to pixel-optimal selection at $r = 32$. Gradient reduction is $100(G_{\text{pixel}} - G_{\text{equal}})/G_{\text{pixel}}$; PSNR decrease is $\text{PSNR}_{\text{pixel}} - \text{PSNR}_{\text{equal}}$

Image	Gradient-error reduction (%)	Raw PSNR decrease (dB)
Astronaut	0.498	0.082
Brick	0.000	0.000
Camera	0.464	0.002
Chelsea	0.817	0.473
Clock	1.698	0.373
Coffee	1.694	0.971
Coins	0.618	0.286
Grass	0.517	0.123
Gravel	0.056	0.007
Moon	0.541	0.014
Motorcycle	0.491	0.060
Rocket	1.375	0.108

The coffee reconstructions in Figure 3 illustrate the effect of these choices. At $B = 96$, pixel-optimal selection retains 29 red, 34 green, and 33 blue components. Equal-weight selection retains 21, 41, and 34, including 13 positions beyond the equal-channel prefix length. PSNR falls from 28.385 to 27.414 dB, and clipped SSIM falls from 0.8040 to 0.7895. The

smaller gradient residual is therefore accompanied by losses in both pixel accuracy and this structural measure.

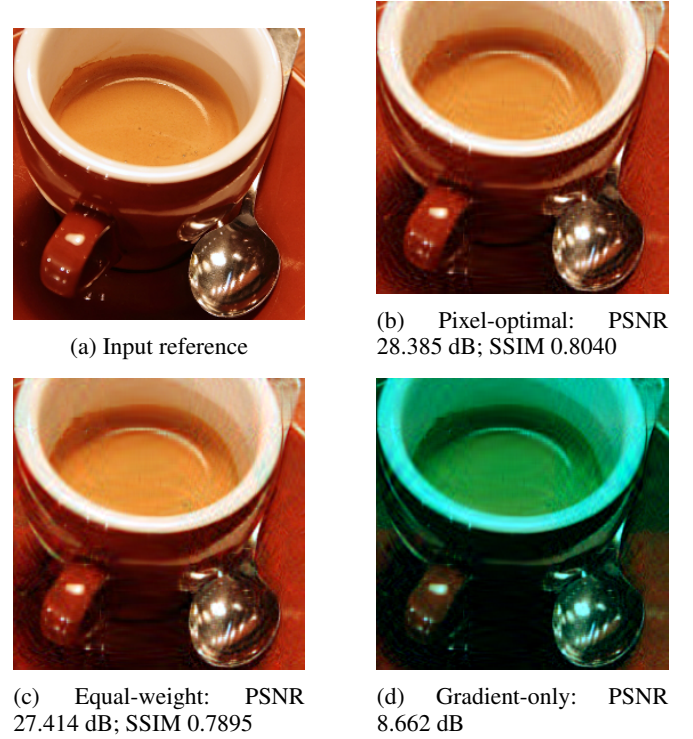


Figure 3: Coffee input and reconstructions at $B = 96$ ($r = 32$). All panels use the same crop and intensity scale, with values clipped to $[0, 1]$ for display. PSNR is measured before clipping; SSIM after clipping. No contrast enhancement or rescaling is applied. Photograph: Rachel Michetti, CC0

Gradient-only selection produces a much larger change. Its normalized pixel residual is 0.542742, compared with 0.005785 for pixel-optimal selection: a factor of 93.819. PSNR falls by 19.723 dB, although gradient error improves by only 1.768%. The displayed reconstruction has a pronounced change in color and broad intensity structure, consistent with discarding components that carry substantial intensity but little gradient energy. The underlying coffee photograph is credited to Rachel Michetti under CC0; all alterations shown here are computed reconstructions.

D. Weight Dependence and Display-Range Clipping

The weight sweep localizes this large change to the last sampled interval (Figure 4). From $\lambda = 0$ to 0.9, coffee's pixel-error fraction rises from 0.005785 to 0.007712. At $\lambda = 1$, it reaches 0.542742, while gradient error decreases only from 0.483678 to 0.483450. Both $\lambda = 0.9$ and 1 allocate 20, 42, and 34 components to the red, green, and blue channels. The failure therefore results from different retained identities, even though the channel counts are unchanged.

Only the seven marked weights were evaluated; the end-point does not establish how every weight near one behaves.

The retained-index file records the exact components selected at each setting.

Clipping changes the comparison further. At $r = 32$, the fractions of out-of-range coffee values are 0.039459 for pixel-optimal, 0.042460 for equal-weight, and 0.214498 for gradient-only selection. Clipping the gradient-only reconstruction increases its gradient-error fraction from 0.483450 to 0.503430. Projection onto $[0, 1]$ cannot increase squared pixel error against a reference in that interval, but the same guarantee does not apply to neighboring-pixel differences.

At $r = 8$, clipping increases the mean gradient-error fraction of gradient-only selection across all twelve crops from 0.681244 to 0.720598. The selection theorem applies to the raw singular expansion; clipping produces a different array whose gradient error must be measured independently.

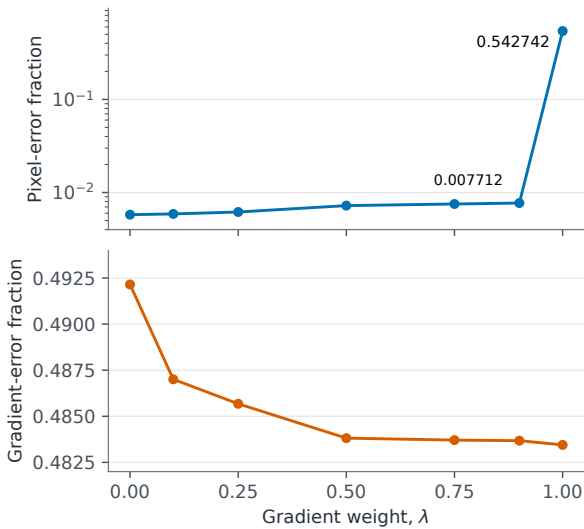


Figure 4: Raw pixel- and gradient-error fractions for the coffee image at $r = 32$ and the seven sampled weights. The pixel-error axis is logarithmic. Connecting lines indicate evaluation order; they do not locate unsampled score crossings

IV. Discussion

A. What the Guarantee Establishes

The additive identity turns subset optimization over a fixed singular dictionary into sorting. The guarantee does not extend to a rank-constrained method that can change the singular vectors, nor does it apply after nonlinear operations such as clipping. Pixel-optimal selection has the smallest pixel residual among subsets of the same size; the corollary shows how the two quadratic errors change as the gradient weight increases. At $r = 32$, the observed gradient gain under equal weighting is less than one percent, whereas the paired PSNR loss varies across images.

B. Brightness and Displayed Images

Constant images are in the null space of the forward-difference operators. A high-energy component with little

gradient contribution can therefore be displaced by a much smaller, more oscillatory component. The two-component example shows that even equal weighting can discard the constant term; the coffee crop shows a pronounced gradient-only failure. A pixel term reduces this effect in that crop but does not guarantee preservation of channel means. Such preservation would require an explicit constraint or a different representation, neither of which was tested.

Clipping a reconstructed image to $[0, 1]$ changes its derivative residual. We therefore calculate PSNR on raw reconstructions and SSIM on clipped ones, and distinguish raw from clipped gradient error. Both the numerical values and the coffee panels support a pixel-fidelity cost; no human preference or semantic-quality study was performed. All color errors refer to the distributed encoded RGB arrays, without a claim of perceptual color uniformity.

C. Limits and Reproducibility

The twelve crops were purposively chosen and are not independent replicates of a larger sample. The 240 reconstructions and additional sweep settings reuse those same photographs. Different scenes, crop positions, color spaces, or resolutions may yield different numerical trade-offs. Full singular decompositions also remain costly for large images, and fewer retained components alone do not establish a smaller encoded file.

The accompanying project records image hashes, component weights, selected indices, raw and clipped measurements, and small-matrix checks. These permit the measured comparisons to be traced to the input arrays and selected dyads. They support the implementation of the exact identities without extending the mathematical claim beyond its stated assumptions.

V. Conclusion

For a fixed singular dictionary, squared pixel and forward-difference residuals decompose over omitted dyads. Sorting a normalized score therefore solves the stated subset problem exactly under a common component budget.

On twelve photographic crops, equal weighting at rank 32 reduces mean normalized gradient error by 0.685% while increasing mean normalized pixel error by 4.155% and reducing paired PSNR by 0.208 dB on average. The coffee example shows a substantially larger pixel loss under gradient-only selection, and clipping changes the measured gradient error. Gradient-aware selection should therefore be assessed alongside pixel fidelity and on the actual displayed reconstruction.

Author Contributions

All authors contributed equally to the conception, development, and preparation of the manuscript. All authors have read and approved the final version of the manuscript for publication.

Conflicts of Interest

The authors declare that there are no conflicts of interest related to this work.

Data Availability

No datasets were generated or analyzed during the present study; therefore, data availability is not applicable.

Funding

This research was supported by operational grants funded by the South Korean Ministry of Science, ICT and Future Planning, including Grant No. 2013R1A1A2057665.

Declaration of Generative AI and AI-Assisted Technologies in the Preparation of the Manuscript

Generative artificial intelligence and AI-assisted tools were used during the preparation and revision of the manuscript. All AI-assisted content was subsequently reviewed and verified by the authors, who take full responsibility for the accuracy, integrity, and final content of the manuscript.

References

- [1] Eckart, C., & Young, G. (1936). The approximation of one matrix by another of lower rank. *Psychometrika*, 1(3), 211–218.
- [2] Ali, M., & Nouy, A. (2020a). Singular value decomposition in Sobolev spaces: Part I. *Zeitschrift für Analysis und ihre Anwendungen*, 39(3), 349–369.
- [3] Ali, M., & Nouy, A. (2020b). Singular value decomposition in Sobolev spaces: Part II. *Zeitschrift für Analysis und ihre Anwendungen*, 39(4), 371–394.
- [4] Halko, N., Martinsson, P.-G., & Tropp, J. A. (2011). Finding structure with randomness: Probabilistic algorithms for constructing approximate matrix decompositions. *SIAM Review*, 53(2), 217–288.
- [5] Gavish, M., & Donoho, D. L. (2014). The optimal hard threshold for singular values is $4/\sqrt{3}$. *IEEE Transactions on Information Theory*, 60(8), 5040–5053.
- [6] Xue, W., Zhang, L., Mou, X., & Bovik, A. C. (2014). Gradient magnitude similarity deviation: A highly efficient perceptual image quality index. *IEEE Transactions on Image Processing*, 23(2), 684–695.
- [7] Wang, Z., Bovik, A. C., Sheikh, H. R., & Simoncelli, E. P. (2004). Image quality assessment: From error visibility to structural similarity. *IEEE Transactions on Image Processing*, 13(4), 600–612.
- [8] Blau, Y., & Michaeli, T. (2018). The perception-distortion tradeoff. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition* (pp. 6228–6237). IEEE.
- [9] Shuman, D. I., Narang, S. K., Frossard, P., Ortega, A., & Vandergheynst, P. (2013). The emerging field of signal processing on graphs: Extending high-dimensional data analysis to networks and other irregular domains. *IEEE Signal Processing Magazine*, 30(3), 83–98.
- [10] van der Walt, S., Schönberger, J. L., Nunez-Iglesias, J., Boulogne, F., Warner, J. D., Yager, N., Gouillart, E., Yu, T., & the scikit-image contributors. (2014). *scikit-image: Image processing in Python*. *PeerJ*, 2, Article e453.
- [11] Scharstein, D., Hirschmüller, H., Kitajima, Y., Krathwohl, G., Nešić, N., Wang, X., & Westling, P. (2014). High-resolution stereo datasets with subpixel-accurate ground truth. In X. Jiang, J. Hornegger, & R. Koch (Eds.), *Pattern recognition* (Lecture Notes in Computer Science, Vol. 8753, pp. 31–42). Springer.
- [12] Harris, C. R., Millman, K. J., van der Walt, S. J., Gommers, R., Virtanen, P., Cournapeau, D., Wieser, E., Taylor, J., Berg, S., Smith, N. J., Kern, R., Picus, M., Hoyer, S., van Kerkwijk, M. H., Brett, M., Haldane, A., Fernández del Río, J., Wiebe, M., Peterson, P., ... Oliphant, T. E. (2020). Array programming with NumPy. *Nature*, 585(7825), 357–362.
- [13] Virtanen, P., Gommers, R., Oliphant, T. E., Haberland, M., Reddy, T., Cournapeau, D., Burovski, E., Peterson, P., Weckesser, W., Bright, J., van der Walt, S. J., Brett, M., Wilson, J., Millman, K. J., Mayorov, N., Nelson, A. R. J., Jones, E., Kern, R., Larson, E., ... SciPy 1.0 Contributors. (2020). *SciPy 1.0: Fundamental algorithms for scientific computing in Python*. *Nature Methods*, 17(3), 261–272.
- [14] Hunter, J. D. (2007). Matplotlib: A 2D graphics environment. *Computing in Science & Engineering*, 9(3), 90–95.

...