

RESEARCH

Open Access



A multi-frame super-resolution method based on the variable-exponent nonlinear diffusion regularizer

Baraka Jacob Maiseli^{1,2*}, Ogada Achieng Elisha³ and Huijun Gao²

Abstract

In this work, the authors have proposed a multi-frame super-resolution method that is based on the diffusion-driven regularization functional. The new regularizer contains a variable exponent that adaptively regulates its diffusion mechanism depending upon the local image features. In smooth regions, the method favors linear isotropic diffusion, which removes noise more effectively and avoids unwanted artifacts (blocking and staircasing). Near edges and contours, diffusion adaptively and significantly diminishes, and since noise is hardly visible in these regions, an image becomes sharper and resolute—with noise being largely reduced in flat regions. Empirical results from both simulated and real experiments demonstrate that our method outperforms some of the state-of-the-art classical methods based on the total variation framework.

Keywords: Super-resolution; Regularization; Image reconstruction; Diffusion

1 Introduction

1.1 Overview of the super-resolution methods

Super-resolution image reconstruction refers to a collection of various image-processing tools used to reconstruct high-resolution images from their corresponding degraded versions ([1–5], and references therein). The term resolution, as used in this work, is the number of pixels per unit area. Therefore, a high-resolution image with both appreciable subjective and objective qualities has a larger pixel count compared with a low-resolution image.

Humans are naturally inclined to vouch for high-quality images. Perhaps the reason for this inclination is the fact that such images contain more information and, hence, are easier to comprehend and interpret. In machine learning and computer vision tasks, detailed images are useful to accurately and robustly highlight and extract critical features, such as edges and contours. These typical examples make super-resolution a vital technology.

Super-resolution methods attempt to address the hardware limitations for improving image qualities. The advantage of these methods is that they discourage hardware modifications to achieve high-resolution images and, thus, are cheap and promote portability. The hardware approach, on the contrary, prompts for changes in the sensor, chip size, and device's internal circuitry to meet the same goal of the super-resolution process. These hardware changes come at the following costs: (1) the current technology disallows further reduction of the sensor's pixel size (which increases pixel count and, therefore, improves resolution) as the process may introduce unwanted shot noise that degrades images; (2) increasing the chip size, which improves the spatial resolution, raises capacitance that slows down the charge transfer rate; and (3) additional weight due to hardware change may limit particular applications, such as remote sensing or unmanned aerial surveillance.

The methods for solving the super-resolution problem can be put into two groups, namely multi-frame and single-frame. The former group uses a single image to generate its corresponding high-resolution version [6–9]. Despite the extensive review, single-frame-based super-resolution approaches—interpolation (nearest neighbor, bilinear, and bicubic) and example-based [10–13]—tend

*Correspondence: barakamaiseli@yahoo.com

¹Department of Electronics and Telecommunication Engineering, College of Information and Communication Technologies, University of Dar es Salaam, P. O. Box 35194, 00255 Dar es Salaam, Tanzania

²Research Institute of Intelligent Control and Systems, Harbin Institute of Technology, 150001 Harbin, China

Full list of author information is available at the end of the article

to generate unpleasing results due to insufficient amount of input images. For example, interpolation methods are known for producing blurry and aliased images. Besides, example-based approaches are computationally inefficient and may, therefore, be unsuitable for real-time processing. A thorough discussion on single-frame super-resolution methods is beyond the scope of our work. However, for a comprehensive review of this category of methods, we refer interested readers to the works in [7, 8, 14–20]

A multi-frame super-resolution approach—which this paper is based upon—attempts to generate a high-resolution image by fusing some pieces of information from multiple degraded images of the same scene [21–24]. The fundamental premise of this approach is that pairs of similar images are likely to differ by rotations, translations, and/or affine transformations due to (1) shaking of the camera, (2) different capturing and exposure times, and (3) relative motion between camera and scene. We can exploit this additional information to produce a more resolute scene by computing the values of these attributes and using them to align the frames onto a common high-resolution grid. The rest of this work discusses the multi-frame super-resolution methods.

The super-resolution problem is highly ill-posed due to insufficient number of the observed low-resolution frames. A classical approach to address this problem is called regularization. Authors have, therefore, proposed different types of priors.

A classical and, perhaps, popular prior is based on the Tikhonov model [25, 26]. This prior introduces into the restoration problem a smoothness constraint that removes extraneous noise from an image. The weakness of the Tikhonov prior is that it tends to destroy edges—an effect that degrades images. Therefore, the prior has captured the interest of many researchers to develop models that simultaneously suppress noise and preserve critical image features: Huber Markov random field (Huber-RMF) [27, 28], edge-adaptive RMF [29, 30], sparse directional [31, 32], and total variation (TV) [33–35].

Of the aforementioned priors, TV has attracted more attention of researchers as it generates results with pleasing objective and subjective qualities. The major weaknesses of the TV model are blocking and staircasing effects, false-edge generation near edges, and non-differentiability property at zero—a situation that makes the numerical implementation rather challenging. The TV model was initially applied in image denoising [36]. Later, the model was adapted to other applications: super-resolution [33], MRI medical image reconstruction [37], inpainting [38], and deblurring [39]. In this work, we explore some classical TV-based approaches to address the super-resolution problem.

In 2008, Marquina et al. proposed a convolutional model for a super-resolution problem based on the constrained TV framework [33]. In their work, the authors introduced the Bregman algorithm as an iterative refinement step to enhance the spatial resolution. The results demonstrate that the method generates detailed and sharper images, but blocking and staircasing artifacts are still evident.

In [35], Farsiu and colleagues proposed a bilateral total variation (BTV) prior, which is based on the L^1 norm minimization and the bilateral filter regularization functional, for a multi-frame super-resolution problem. Their method is computationally inexpensive and robust against errors caused by motion and blur estimations and generates images that are convincingly sharper. However, bilateral filters that the method derives its strengths from are known to introduce artifacts like staircasing and gradient reversal. Additionally, the BTV inadequately addresses the partial smoothness of an image [40]. Besides, the numerical implementation of the L^1 norm component is challenging as it masks the super-resolution data term.

In [34], Ng et al. applied the TV prior to address the following issues in the super-resolution video reconstruction: noise, blurring, missing regions, compression artifacts, and motion estimation errors. The authors demonstrated the efficacy of their method in several cases of motions and degradations and provided the experimental results that outperform some other classical super-resolution methods.

Ren et al. proposed a super-resolution method, which is based on the fractional order TV regularization, with a focus to handle fine details, such as textures, in the image [41]. Results show that their approach addresses to some extent the weaknesses of the traditional TV.

In [21], Li et al. attempted to address the drawbacks of the global TV by proposing two regularizing functionals, namely locally adaptive TV and consistency of gradients, to ensure that edges are sharper and flat regions are smoother. The method heavily depends on the gradient details of an image, a feature that may produce pseudo-edges in noisy homogeneous regions. Note that both noise and edges are image features with high-gradient (or high-intensity) values. As Li's method is gradient-dependent, it may equally treat both noise and edges, and this may generate unwanted artifacts.

Yuan et al. proposed a spatially weighted TV model for the multi-frame super-resolution problem [42]. Their model incorporates a spatial information indicator (difference curvature) that locally identifies the spatial properties of the individual pixels, thus providing the necessary level of regularization. The authors employed the majorization-minimization algorithm to optimize their formulation. Results show that the Yuan et al. method overcomes some challenges of the original TV model (discourages piecewise constant solutions and is less sensitive

to regularization parameters) and outperforms the Li et al. method. But under severe noise conditions, the Yuan et al. approach fails as it is pixel-unit-based [40].

Recently, Weili et al. proposed an adaptive TV-driven super-resolution method that provides convincing results compared with those generated by the standard TV model [40]. The authors incorporated into the classical super-resolution formulation a feature-sensitive prior that approximates the L^1 norm near edges, and this helps to efficiently highlight these critical features. In flat regions, the prior approximates the L^2 norm, thus providing noise removal.

Inspired by the weaknesses of the standard TV and its variants, this work proposes an alternative framework based on nonlinear diffusion processes—which have yielded promising results in image-denoising applications [36, 43–47]—to address the super-resolution ill-posedness. Our diffusion-driven prior includes an adaptive kernel that sensitively and dynamically updates its value in accordance with the scanned local image features. That is, the kernel is linear isotropic in flat regions and nonlinear anisotropic near edges. Being locally adaptive, the model generates more resolute scenes and avoids blocking artifacts inherent in the conventional TV.

1.2 A classical multi-frame image degradation model

Imaging devices, such as cameras, endure unavoidable hardware limitations and influence of external noise. Therefore, an image signal that passes through the imaging system usually undergoes degradations before reaching its destination (Fig. 1).

Let us consider that u is the unknown high-resolution image. In the image degradation model (Fig. 1), u is first

acted upon by the warping operator, W_k , which rotates and translates it. Then, the degraded version of u proceeds to the next stage, where it is blurred by the blurring operator, B_k . In the next stage, the dimension (width $\times$ height) of the degraded (warped and blurred) u is decreased by the decimation operator, D_k . Finally, noise (assumed to be additive white Gaussian), η_k , is added as a further image degradation agent. Consequently, a sequence of low-resolution images, $\{y_k\}$, where $k = 1, 2, 3, \dots, M$ is generated. As the model is multi-frame-based, the degradation process incorporates a sequence of composite operators— $W_1 B_1 D_1, \dots, W_k B_k D_k, \dots, W_M B_M D_M$ —which are applied to u to produce $y_1, \dots, y_k, \dots, y_M$, respectively.

Therefore, we intuitively represent the multi-frame image degradation model by the equation

$$y_k = W_k B_k D_k u + \eta_k \quad (1)$$

Treating η_k as energy, our objective is to minimize it using the L^2 norm that is known for its ability to suppress noise. Therefore, the minimization problem becomes

$$\min_u \left\{ \frac{1}{2M} \sum_{k=1}^M \|y_k - W_k B_k D_k u\|_{L^2(\Omega)}^2 \right\}, \quad (2)$$

where Ω is the supporting domain of u . Applying the Euler-Lagrange optimization approach, and embedding the resulting solution into a dynamical system, we get

$$\frac{\partial u}{\partial t} = \frac{1}{M} \sum_{k=1}^M W_k' B_k' D_k' (W_k B_k D_k u - y_k). \quad (3)$$

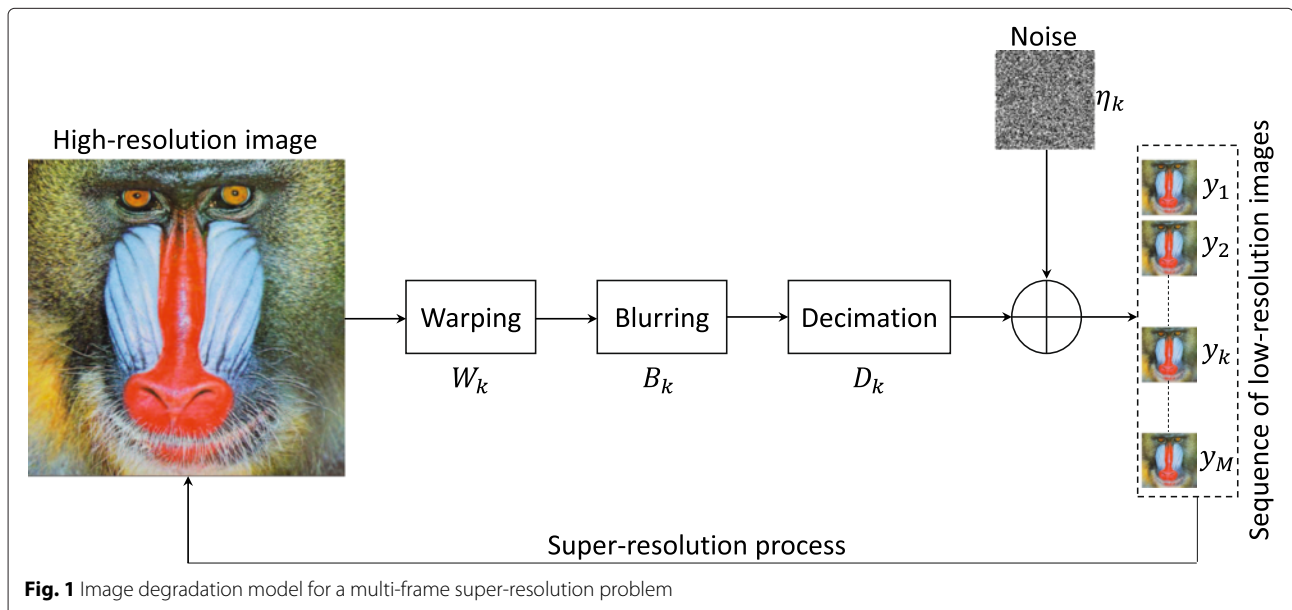


Fig. 1 Image degradation model for a multi-frame super-resolution problem

2 Multi-frame super-resolution process

2.1 Background

The fundamental aim of the super-resolution process is to improve the spatial resolution of an image while suppressing unwanted artifacts. A well-known approach to achieve this goal is to incorporate a regularizing potential into the super-resolution energy functional. Therefore, we can reformulate (2) as a variational problem

$$\min_{u \in BV \cap L^2(\Omega)} \left\{ \frac{1}{2M} \int_{\Omega} \sum_{k=1}^M (y_k - W_k B_k D_k u)^2 dx + \int_{\Omega} \psi(|\nabla u|) |\nabla u| dx + \frac{\lambda}{2} \int_{\Omega} (u - f)^2 dx \right\}, \quad (4)$$

where the terms from left are defined as follows: super-resolution or data, regularization potential, and fidelity, respectively. From (4), the coefficient function, $\psi(s)$, simultaneously detects edges and penalizes the norm of the image gradient, and the fidelity term, which contains λ as the regularizing parameter, establishes a trade-off between the evolving image, u , and the initial guess, f .

The parameterized (time-dependent) Euler-Lagrange form of (4) is

$$\begin{aligned} \frac{\partial u}{\partial t} = & \frac{1}{M} \sum_{k=1}^M W'_k B'_k D'_k (W_k B_k D_k u - y_k) \\ & + \operatorname{div} \left(\left(\frac{\psi(|\nabla u|)}{|\nabla u|} + \psi'(|\nabla u|) \right) \nabla u \right) \\ & - \lambda(u - f). \end{aligned} \quad (5)$$

In [36], Rudin et al. proposed $\psi(s) = 1$, and plugging this value into (5) yields

$$\begin{aligned} \frac{\partial u}{\partial t} = & \frac{1}{M} \sum_{k=1}^M W'_k B'_k D'_k (W_k B_k D_k u - y_k) + \operatorname{div} \left(\frac{1}{|\nabla u|} \nabla u \right) \\ & - \lambda(u - f), \end{aligned} \quad (6)$$

which is the multi-frame super-resolution model based on the Rudin-Osher-Fatemi (ROF) regularizing functional. Although the ROF model suppresses noise and effectively recovers edges, it has some limitations, as noted by Ogada and colleagues [43]. Firstly, the formulation favors piecewise-constant solutions that result into staircasing effects or even generation of false edges. Secondly, the ROF model tends to reduce contrast in homogeneous or noise-free image regions. Thirdly, the TV diffusion model—despite its anisotropic property—produces a process that only diffuses in a direction that is orthogonal to the gradients of the image contours. This has a consequence of producing blockiness in the results. And lastly, the ROF evolution system contains a $\frac{1}{|\nabla u|}$ component that runs into a spike when $|\nabla u| = 0$ in flat regions. To

solve this challenge, the numerical implementations of the model usually incorporate a lifting parameter $0 < \epsilon < 1$, where ϵ is made too small. That is, the diffusion coefficient in the divergence part of (6) is replaced by $\frac{1}{|\nabla u| + \epsilon}$. In our view, this modification limits the accuracy of the results and may even cause instabilities.

Motivated by the weaknesses of the ROF model, Ogada et al. proposed an alternative method for image denoising that focuses on selecting an appropriate value of $\psi(s)$. Their approach provides a criteria for deciding the nature of $\psi(s)$ in terms of linearity, sub-linearity, and super-linearity. This ensures that the resulting regularizing functional is strictly convex and grows linearly. In image-denoising problems, linear growth and convex functionals are known to generate appealing results. Additionally, the authors' diffusion equation contains the denominator that never collapses to zero even in smooth regions (where $|\nabla u| = 0$), thus avoiding the singularity risk like that encountered in the TV problem.

Therefore, Ogada et al. selected a sub-linear growth functional coefficient

$$\psi(|\nabla u|) = \frac{|\nabla u|}{1 + |\nabla u|}. \quad (7)$$

Plugging $\psi(s)$ from (7) into (5) gives

$$\begin{aligned} \frac{\partial u}{\partial t} = & \frac{1}{M} \sum_{k=1}^M W'_k B'_k D'_k (W_k B_k D_k u - y_k) + \operatorname{div} \left(\frac{2 + |\nabla u|}{(1 + |\nabla u|)^2} \nabla u \right) \\ & - \lambda(u - f). \end{aligned} \quad (8)$$

Decomposing (7) into u_{TT} (tangential) and u_{NN} (normal) components yields

$$\begin{aligned} \frac{\partial u}{\partial t} = & \frac{1}{M} \sum_{k=1}^M W'_k B'_k D'_k (W_k B_k D_k u - y_k) + \left(\frac{2 + |\nabla u|}{(1 + |\nabla u|)^2} \right) u_{TT} \\ & + \left(\frac{2}{(1 + |\nabla u|)^3} \right) u_{NN} - \lambda(u - f), \end{aligned} \quad (9)$$

which possesses the diffusion mechanisms in both tangential and normal directions to the isophote lines. Furthermore, Eq. 9 promotes varying degrees of diffusions depending upon the local image structures, particularly edges and contours.

Therefore, in order to reap the benefits of the model by Ogada et al. and, indeed, the edge-preserving capability and sensitivity of the model to the finer local image structures, we have exponentiated the denominator of the diffusivity of (8) by an adaptive term. Our goal was to make the smoothing functional adapt several models: isotropic diffusion, ROF, and Ogada et al., and to ensure global minimum energy that guarantees uniqueness in the results.

2.2 Proposed model

2.2.1 Model formulation

The smoothing functional proposed by Ogada et al. [43] produces superior results compared with the Perona-Malik (PM) [44], D- α -PM [45], and total variation [36]. The model has been used for denoising applications. In this work, we have modified the model by integrating into its diffusivity an edge-probing variable exponent that is robust against noise. Furthermore, the modified formulation was encapsulated into the classical multi-frame super-resolution model.

Now, we propose a non-variational problem

$$0 = \frac{1}{M} \sum_{k=1}^M W'_k B'_k D'_k (W_k B_k D_k u - y_k) + \operatorname{div} \left(\frac{2 + \frac{|\nabla u|}{\beta}}{\left(1 + \frac{|\nabla u|}{\beta}\right)^{\alpha(x)}} \nabla u \right) - \lambda(u - f), \quad (10)$$

where β is a shape-defining parameter and

$$\alpha(x) = 2 - \frac{1}{1 + \kappa |\nabla(G_\sigma * f)|^2} \quad (11)$$

is an adaptive feature-dependent variable exponent with $\sigma > 0$ and $G_\sigma(x) = \frac{1}{4\pi\sigma} \exp\left(\frac{-|x|^2}{4\sigma^2}\right)$ (Fig. 2). The authors in [46] found that $\sigma = 0.50$ and $0.0025 < \kappa < 0.025$ produced promising restoration results, and indeed, their findings worked well for our case. Equation 10 is usually solved by embedding it into a dynamical system, which

is then evolved until steady state conditions are attained. Therefore, parameterizing the equation in time yields

$$\begin{aligned} \frac{\partial u}{\partial t} &= \frac{1}{M} \sum_{k=1}^M W'_k B'_k D'_k (W_k B_k D_k u - y_k) + \operatorname{div} \left(\frac{2 + \frac{|\nabla u|}{\beta}}{\left(1 + \frac{|\nabla u|}{\beta}\right)^{\alpha(x)}} \nabla u \right) \\ &\quad - \lambda(u - f) \quad \text{for } (x, t) \in \Omega \times (0, T), \\ u(x, 0) &= f, \quad x \in \Omega, \\ \frac{\partial u}{\partial n} &= 0, \quad (x, t) \in \partial\Omega \times (0, T), \end{aligned} \quad (12)$$

which can be implemented in the computer using the appropriate numerical schemes. In this paper, we used the four-point neighborhood explicit scheme to implement (12), as detailed in the later sections.

Figure 3 shows the block diagram that implements our model in (12). From the illustration, a sequence of M frames, $y_1, y_2, \dots, y_k, \dots, y_M$, is first captured by an imaging device. Then, the values of the motion parameters (rotations and translations) between $\{y_{k+1}\}_{k=1,2,\dots,M}$ and y_1 (reference frame) are computed using the Keren et al. algorithm—which is chosen for its reliability, robustness, accuracy, and computational efficiency. Next, the frames are aligned using the computed motion values and projected onto the high-resolution grid. Next, the Zomet et al. method is used to robustly detect outliers in the data. The actual reconstruction is done in the following step, and we used the steepest descent approach for this purpose. Lastly, the proposed regularizing functional is incorporated to address the super-resolution ill-posedness and also to address noise issues in the image. The program's iteration exit point is determined using the L^2 error norm

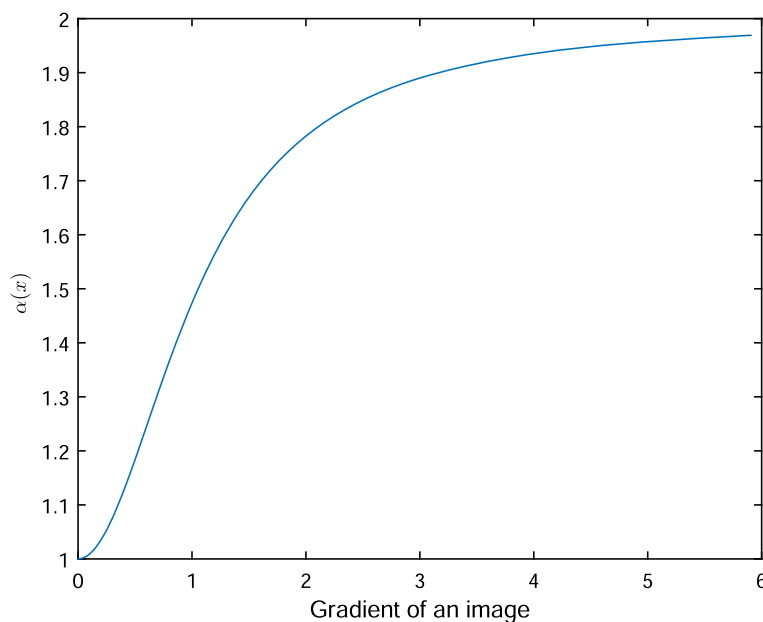
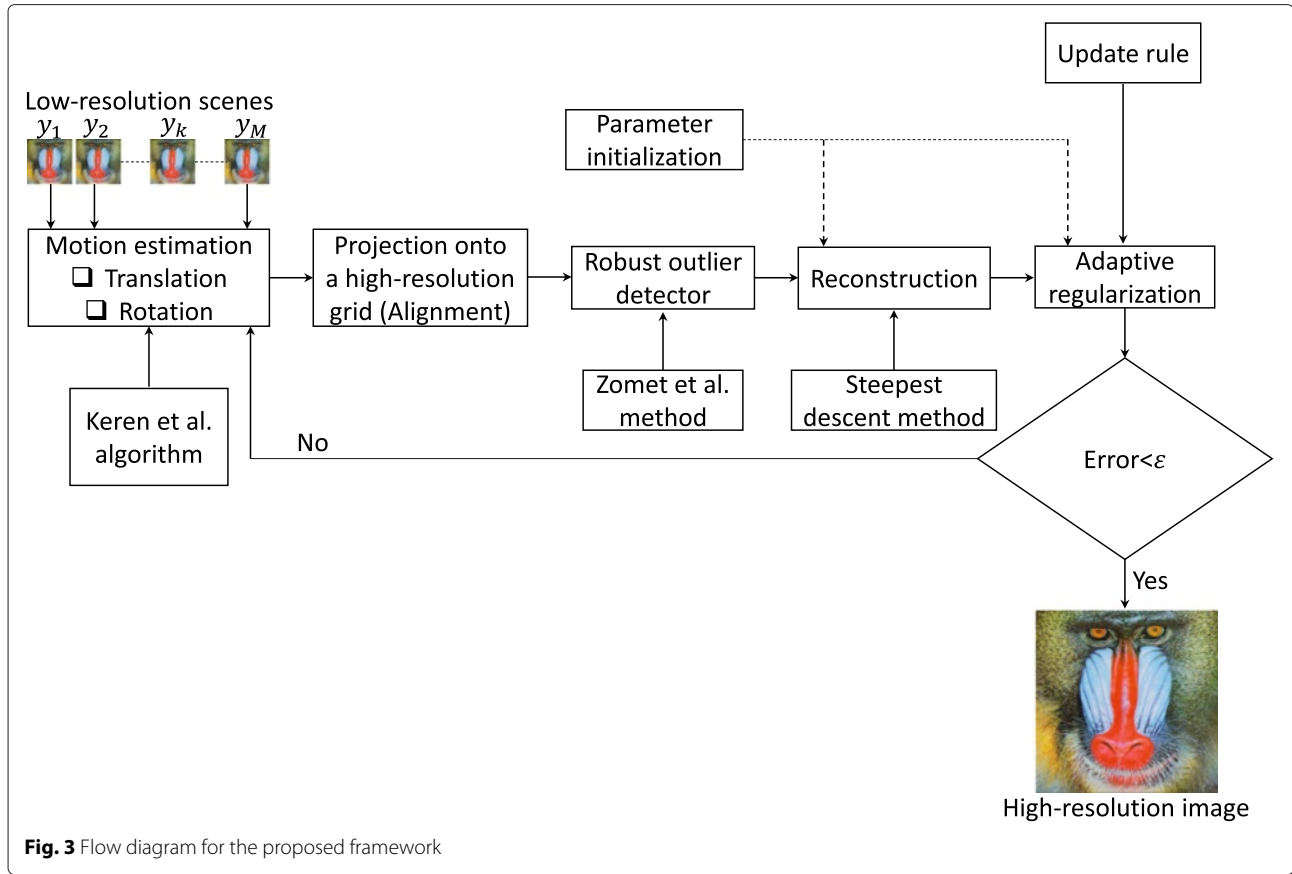


Fig. 2 Variation of the feature-dependent variable exponent, $\alpha(x)$, with respect to the image gradient



$$\frac{\|u^{(n)} - u^{(n-1)}\|_2^2}{\|u^{(n)}\|_2^2} < \varepsilon, \quad (13)$$

where $0 < \varepsilon < 1$ is a tuning constant that determines the final error of the results.

2.2.2 Physical significance and roles of $\alpha(x)$

The variable, $\alpha \in [1, 2]$, sweeps values between one and two according to the local features of an image (Fig. 4). From (11), we observe the following cases:

1. In flat regions ($|\nabla(G_\sigma * f)| \rightarrow 0$), $\alpha = 1$. Substituting this value of α into (12) yields

$$\begin{aligned} \frac{\partial u}{\partial t} = & \frac{1}{M} \sum_{k=1}^M W'_k B'_k D'_k (W_k B_k D_k u - y_k) \\ & + \operatorname{div} \left(\frac{2 + \frac{|\nabla u|}{\beta}}{1 + \frac{|\nabla u|}{\beta}} \nabla u \right) \\ & - \lambda(u - f). \end{aligned}$$

Now, expanding the divergence part of this equation produces

$$\begin{aligned} \frac{\partial u}{\partial t} = & \frac{1}{M} \sum_{k=1}^M W'_k B'_k D'_k (W_k B_k D_k u - y_k) \\ & + \operatorname{div} \left(\frac{1}{1 + \frac{|\nabla u|}{\beta}} \nabla u \right) + \Delta u, \end{aligned} \quad (14)$$

which combines two regularizing models, namely ROF and isotropic diffusion. Therefore, the formulation can isotropically remove noise in these (flat) regions. Additionally, if β is carefully tuned, we may preserve weak edges due to the presence of a regularizing component—middle term of (14)—that is similar to that of TV.

2. Near edges ($|\nabla(G_\sigma * f)| \rightarrow \infty$), $\alpha = 2$. Thus, Eq. 12 becomes

$$\begin{aligned} \frac{\partial u}{\partial t} = & \frac{1}{M} \sum_{k=1}^M W'_k B'_k D'_k (W_k B_k D_k u - y_k) \\ & + \operatorname{div} \left(\frac{2 + \frac{|\nabla u|}{\beta}}{\left(1 + \frac{|\nabla u|}{\beta}\right)^2} \nabla u \right) \\ & - \lambda(u - f), \end{aligned}$$

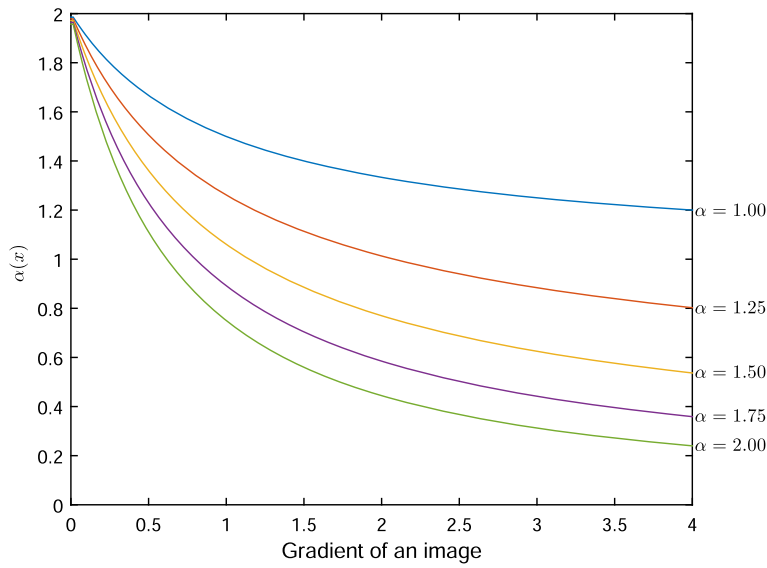


Fig. 4 Variation of the proposed diffusivity with respect to different values of $\alpha(x)$

which contains a regularizing part proposed by Ogada et al. [43]. This formulation helps to suppress noise, avoid diffusion of edges, and enhance the spatial resolution of an image.

- From the above two cases, I and II, we see that $\alpha(x)$ plays another role of segmenting an image into two subregions, namely Ω_1 (flat regions; $\alpha = 1$) and Ω_2 (edges and contours; $\alpha = 2$). Hence, α is an edge-defining variable. An important aspect of α is that it contains a convolution operator between the Gaussian kernel and the image, $G_\sigma * f$, which helps to suppress noise and other unwanted artifacts—thus detecting useful image features robustly, even under harsh imaging conditions (Figs. 5 and 6).

2.2.3 Properties of the model

To comprehensively understand the (mechanical) properties of our model, consider two orthogonal vectors, $T(x) = (u_x, u_y)/|\nabla u|$ and $N(x) = (-u_y, u_x)/|\nabla u|$, in an image, u , where u_x and u_y are the first-order partial derivatives of u in the x - and y -direction, respectively, and $|\nabla u| \neq 0$. Additionally, let u_{TT} (tangential) and u_{NN} (normal) be the second-order partial derivatives of u , representing diffusions, in the directions of T and N , respectively. Defining u_{TT} and u_{NN} as

$$u_{TT} = T' \nabla^2 u T = \frac{1}{|\nabla u|^2} (u_x^2 u_{yy} + u_y^2 u_{xx} - 2u_x u_y u_{xy}) \text{ and}$$

$$u_{NN} = N' \nabla^2 u N = \frac{1}{|\nabla u|^2} (u_x^2 u_{xx} + u_y^2 u_{yy} + 2u_x u_y u_{xy}),$$

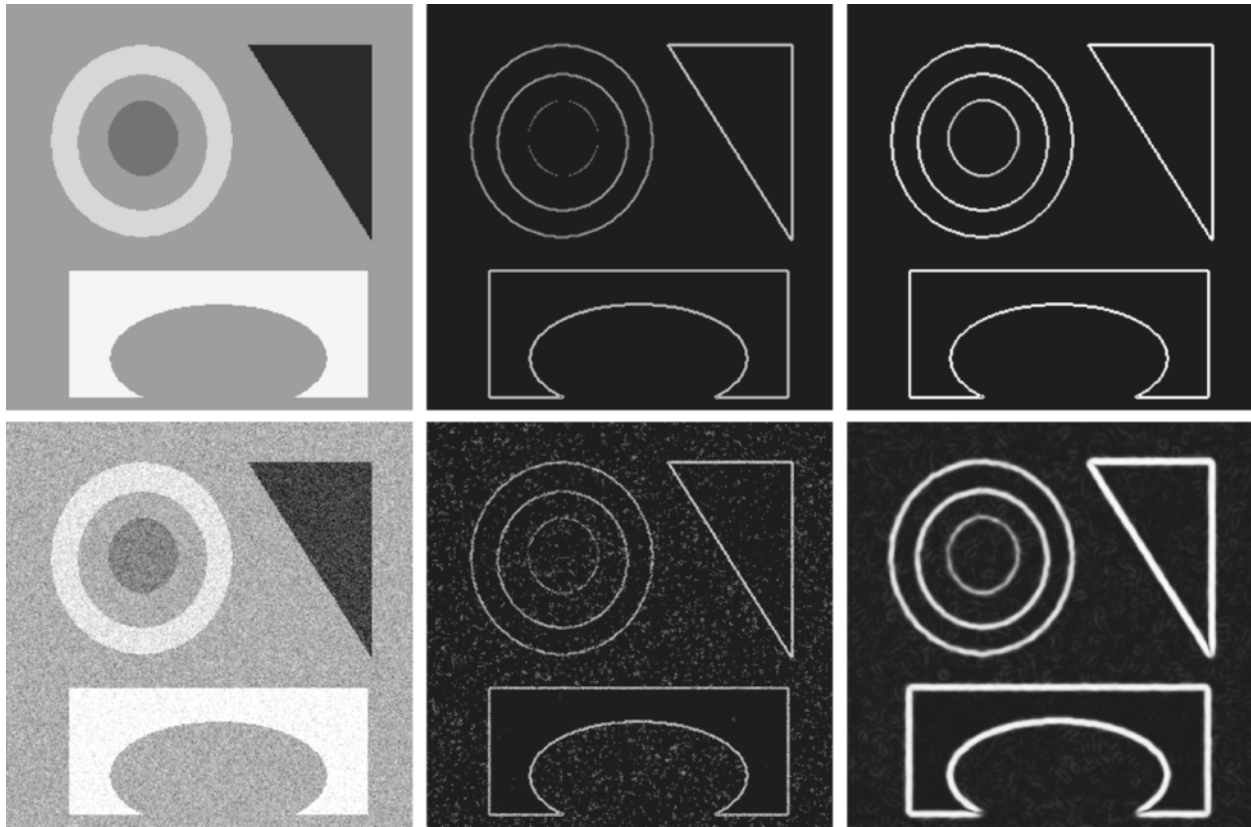
where $T' = (1/|\nabla u|) \begin{pmatrix} u_y \\ u_x \end{pmatrix}$, $N' = (1/|\nabla u|) \begin{pmatrix} u_x \\ -u_y \end{pmatrix}$, and $\nabla^2 u = \begin{pmatrix} u_{xx} & u_{xy} \\ u_{xy} & u_{yy} \end{pmatrix}$ is the Hessian matrix; Eq. 12 can compactly be re-written in terms of u_{TT} and u_{NN} as

$$\begin{aligned} \frac{\partial u}{\partial t} = & \frac{1}{M} \sum_{k=1}^M W_k' B_k' D_k' (W_k B_k D_k u - y_k) + \left(\frac{2 + \frac{|\nabla u|}{\beta}}{\left(1 + \frac{|\nabla u|}{\beta}\right)^{\alpha(x)}} \right) u_{TT} \\ & + \left(\frac{\left(1 + \frac{|\nabla u|}{\beta}\right) \left(\left(2 + \frac{|\nabla u|}{\beta}\right) + 1 \right) - \left(2 + \frac{|\nabla u|}{\beta}\right) \alpha(x)}{\left(1 + \frac{|\nabla u|}{\beta}\right)^{\alpha(x)+1}} \right) u_{NN} - \lambda(u - f). \end{aligned} \quad (15)$$

Equation 15 possesses some interesting properties worth noting: in flat regions ($|\nabla u| \rightarrow 0$ and $\alpha \rightarrow 1$), the equation reduces to

$$\frac{\partial u}{\partial t} = \frac{1}{M} \sum_{k=1}^M W_k' B_k' D_k' (W_k B_k D_k u - y_k) + (C u_{TT} + u_{NN}) - \lambda(u - f), \quad (16)$$

(C is a constant value) which contains an isotropic diffusion component ($\Delta u = u_{TT} + u_{NN}$ is the pure heat equation) that removes noise uniformly over the regions. Also, near edges ($|\nabla u| \rightarrow \infty$ and $\alpha \rightarrow 2$), the coefficient of u_{NN} —which contains the denominator larger than that of u_{TT} —vanishes faster. Consequently,



(a) Original (up) and noisy (down; additive White Gaussian Noise, $\sigma = 30$)

(b) Ogada et al.

(c) Our method

Fig. 5 Edge maps of a synthetic image: **a** original and noisy (white Gaussian noise) image, **b** image generated from the Ogada et al. model, and **c** image generated using our method

the u_{TT} component that is responsible to preserve edges dominates.

2.3 Numerical implementation

The desire to implement our method using an explicit scheme is attributed to the following reasons: computational efficiency, ability to produce more accurate and appealing results, intuitiveness to be understood and analyzed mathematically, and stability over the time interval defined by the Courant-Friedrichs-Lewy criterion ($0 < \tau \leq 0.25$) [48]. A significant drawback, however, of explicit schemes is that they are susceptible to instabilities for larger iteration steps.

Now, consider a five-point discretized space with four directions: north, south, east, and west (Fig. 7). The gradients of u_{ij} along these directions are, respectively, $\Delta_{ij}^N = u_{i,j+1} - u_{ij}$, $\Delta_{ij}^S = u_{i,j-1} - u_{ij}$, $\Delta_{ij}^E = u_{i+1,j} - u_{ij}$, and $\Delta_{ij}^W = u_{i-1,j} - u_{ij}$,

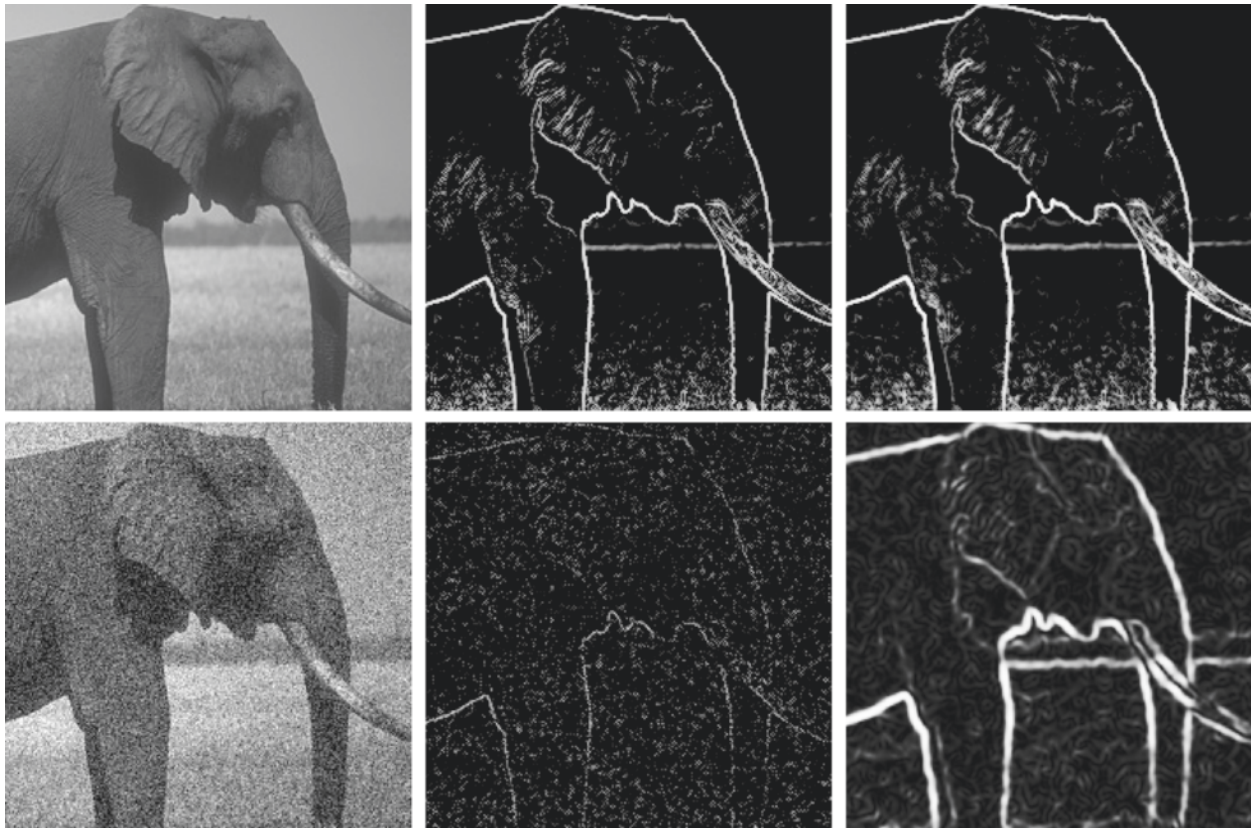
and the corresponding discrete conduction coefficients from the divergence part of (12) are

$$cN_{ij} = \left(\frac{2 + \frac{|\Delta_{ij}^N|}{\beta}}{\left(1 + \frac{|\Delta_{ij}^N|}{\beta}\right)^{\alpha_{ij}}} \right), cS_{ij} = \left(\frac{2 + \frac{|\Delta_{ij}^S|}{\beta}}{\left(1 + \frac{|\Delta_{ij}^S|}{\beta}\right)^{\alpha_{ij}}} \right),$$

$$cE_{ij} = \left(\frac{2 + \frac{|\Delta_{ij}^E|}{\beta}}{\left(1 + \frac{|\Delta_{ij}^E|}{\beta}\right)^{\alpha_{ij}}} \right), \text{ and } cW_{ij} = \left(\frac{2 + \frac{|\Delta_{ij}^W|}{\beta}}{\left(1 + \frac{|\Delta_{ij}^W|}{\beta}\right)^{\alpha_{ij}}} \right),$$

where $0 \leq i \leq P$ and $0 \leq j \leq Q$; P and Q are, respectively, the horizontal and vertical dimensions of u_{ij} . Following the explicit scheme, we discretize the divergence component of our formulation as

$$\text{div}_{ij} = cN_{ij}\Delta_{ij}^N + cS_{ij}\Delta_{ij}^S + cE_{ij}\Delta_{ij}^E + cW_{ij}\Delta_{ij}^W.$$



(a) Original (up) and noisy (down; additive White Gaussian Noise, $\sigma = 30$)

(b) Ogada et al.

(c) Our method

Fig. 6 Edge maps of an elephant image: **a** original and noisy (white Gaussian noise) image, **b** image generated from the Ogada et al. model, and **c** image generated using our method

From the given discretizations, therefore, the steepest descent implementation of the regularized super-resolution method in (12) becomes

$$u_{ij}^{(n+1)} = u_{ij}^{(n)} - \tau \left(R(u_{ij}^{(n)}) - M\mu \text{div}_{ij}^{(n)} + M\mu\lambda (u_{ij}^{(n)} - f_{ij}) \right), \quad (17)$$

where μ is a tuning constant and

$$R(u_{ij}^{(n)}) = \frac{1}{M} \sum_{k=1}^M W'_k B'_k D'_k (W_k B_k D_k u_{ij}^{(n)} - (y_k)_{ij}) \quad (18)$$

is a super-resolution result at the n th iteration. The boundary conditions of (17) are as follows: $u_{ij}^{(0)} = (y_{ij})_1 = y_1(ih, jh)$, $u_{i,0}^{(n)} = u_{i,1}^{(n)}$, $u_{0,j}^{(n)} = u_{1,j}^{(n)}$, $u_{p,j}^{(n)} = u_{p-1,j}^{(n)}$, and $u_{i,Q}^{(n)} = u_{i,Q-1}^{(n)}$

2.4 Performance evaluation of the super-resolution models

To quantitatively evaluate the performance of different super-resolution models, we have used three metrics: peak-signal-to-noise ratio (PSNR) [49, 50], edge similarity (ESIM) [45], and mean structure similarity (MSSIM) [51]. The PSNR measures signal strength relative to noise in the image and is defined by the equation

$$\text{PSNR} = 10 \log_{10} \left(\frac{255^2 PQ}{\|u - f\|_2^2} \right). \quad (19)$$

Though widely applied in many image-processing disciplines, PSNR fails to explain the quality of edges in the restored scenes. Guo et al. addressed this limitation by proposing the metric

$$\text{ESIM} = \widetilde{\text{PSNR}}(\text{EM}(u), \text{EM}(f)) \quad [45], \quad (20)$$

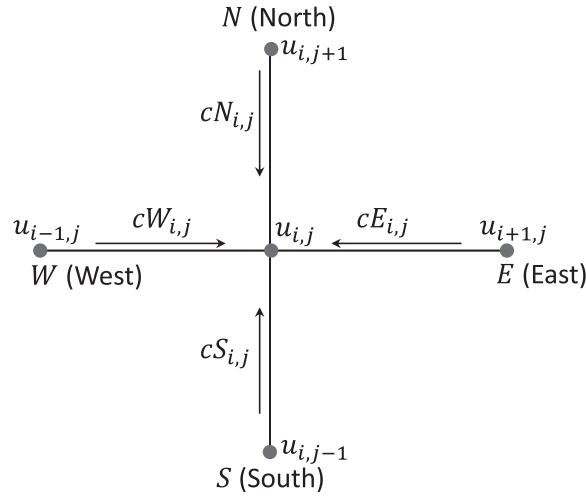


Fig. 7 Explicit numerical scheme

where $EM(u)$ and $EM(f)$ are, respectively, the edge maps of u (restored image) and f (initial ideal image), and

$$\widetilde{PSNR} = 10 \log_{10} \frac{PQ |\max(f) - \min(f)|^2}{\|u - f\|_2^2}$$

is as proposed in [52]. Moreover, both PSNR and ESIM inadequately address the perceptual qualities of scenes. Wang and colleagues proposed a metric (MSSIM) that explains the statistical inter-dependency of pixels in scenes, thus quantifying their visual appeals [51]. The MSSIM is governed by the equation

$$MSSIM = \frac{(2\mu_u\mu_f + c_1)(2\sigma_{uf} + c_2)}{(\mu_u^2 + \mu_f^2 + c_1)(\sigma_u^2 + \sigma_f^2 + c_2)}, \quad (21)$$

where the variables, respectively defined for u and f are as follows: μ_u and μ_f , mean; σ_u^2 and σ_f^2 , variance; and σ_{uf} , covariance. And c_1 and c_2 are stabilizing constants.

3 Experiments

We conducted experiments from both simulated and real environments to test and compare the performance of the proposed method and some classical super-resolution methods, namely (Total variation) [33], (Spatially weighted total variation) [42], (Fractional order total variation) [41], and (Adaptive total variation) [40]. In the first experiment, we degraded each of the original high-resolution images of Boat, Bridge, Building, Fish, Goldhill, Lena, Mandrill, and Wheel (<http://decsai.ugr.es/cvg/CG/base.htm>) (Fig. 8) to generate the corresponding sequence of ten low-resolution images. Next, we applied TV-SR, SWTV-SR, FTV-SR, ATV-SR, and the new method to restore the respective original versions of the degraded image sequences (Fig. 9). These procedures were repeated in the second experiment but with the input

images being the degraded low-resolution video frames of EIA (<https://users.soe.ucsc.edu/~milanfar/software/sr-images.html>) (Fig. 10). The last experiment tested the efficacy of our method for two conditions of the variable exponent, $\alpha(x)$: fixed ($\alpha = 1$ and 2) and adaptive (Fig. 11).

4 Results and discussions

The visual results demonstrate that the proposed method outperforms in several cases compared with some state-of-the-art classical methods (Figs. 9 and 10). Both simulated and real experiments show that the new approach generates appealing images that are sharper and detailed. Other methods, such as TV-SR, SWTV-SR, and FTV-SR, reveal obvious artifacts—ringing, blocking, and staircasing (Fig. 9). The last experiment proves that the adaptive nature of our formulation helps to preserve useful image features and suppress noise; setting constant the edge-locating variable exponent, $\alpha(x)$, in the proposed model lowers its performance (Fig. 11). We observed earlier in Section 2.2.2—“Physical significance and roles of $\alpha(x)$ ”—that fixing α to 1 ($\alpha = 1$), for example, makes the super-resolution problem regularized by both TV and isotropic diffusion, and this promotes edge recovery and noise removal, respectively. The results in the third experiment prove this mathematical intuition but also reveal some artifacts that are probably due to fixing $\alpha(x)$, as depicted by Fig. 11c. In Fig. 11e, we observe that adaptively updating $\alpha(x)$ makes the results appear attractive.

Quantitative results further confirm that the proposed model is superior as it achieves higher quality values. Table 1 shows that the new method achieves promising PSNR values in most cases, which implies that the method retains higher signal contents in the restored results. In terms of the edge recovery capabilities, our method performs better as it shows larger values of ESIM

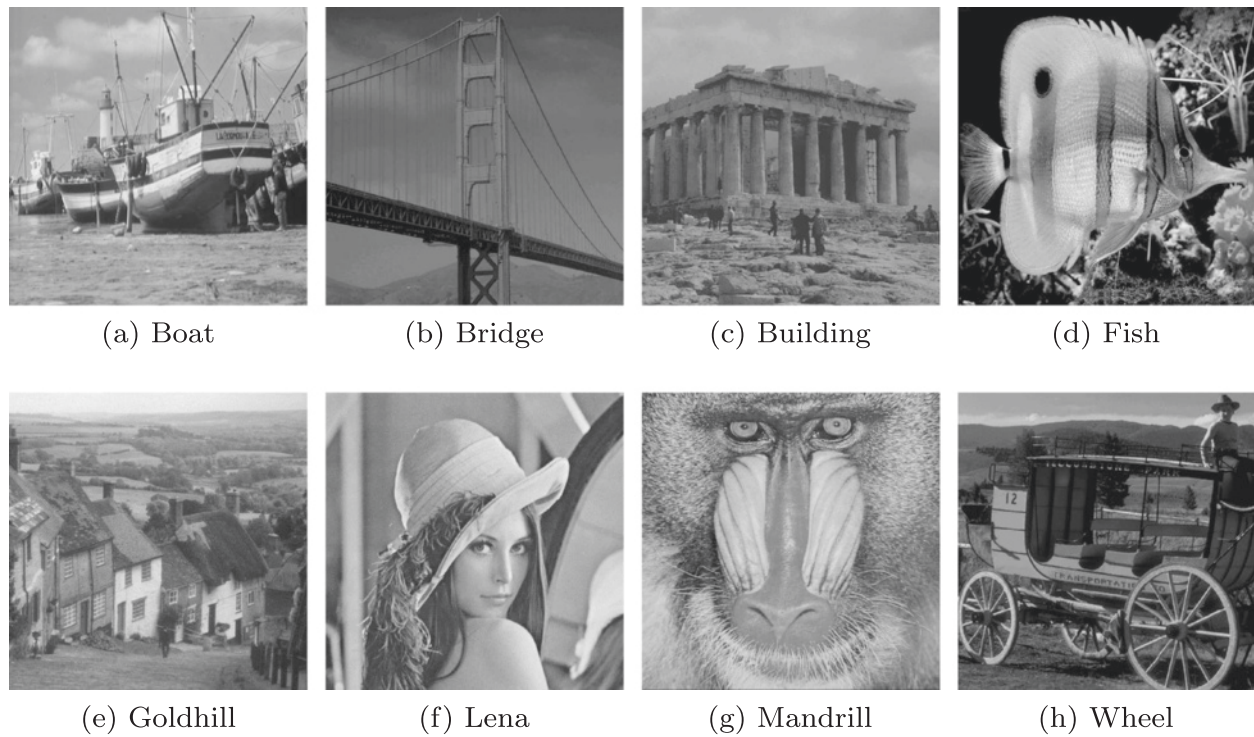


Fig. 8 Original high-resolution images: **a** Boat, **b** Bridge, **c** Building, **d** Fish, **e** Goldhill, **f** Lena, **g** Mandrill, and **h** Wheel

compared with those from the other methods (Table 2). This finding concurs with the subjective assessments that the proposed method generates images with sharper and stronger edges, which raises the value of the ESIM. In Table 3, the results demonstrate that the new method produces MSSIM values that are larger than those produced by TV-SR, SWTv-SR, FTV-SR, and ATV-SR. This quantitative quality assessment justifies the observation that our results are visually appealing (Figs. 9 and 10). Recall that MSSIM favors the human visual system. Furthermore, Table 4 demonstrates that our method offers lower computational times—thus making it a suitable choice in tasks that demands parallel computing.

The performance of ATV-SR closely follows that of our method. Subjectively, the results of these two methods are hard to distinguish (Fig. 9; images of Building and Lena; and Fig. 10). However, numerical results indicate that the proposed approach outperforms in several cases (Tables 1, 2, and 3). For the “Wheel” image, however, Table 1 shows that the ATV-SR outperforms our method by a small amount. Also, the ATV-SR is slightly faster than the proposed method (Table 4), but the deviation is too small that we may assume the two methods perform equally.

Despite the promising performance, the proposed method suffers from one weakness—it tends to slightly blur the output images. This is probably caused by the low-pass filtering operation of the regularizing functional

or inappropriate estimation of the blur function. More research is thus needed to address the limitation. It is worth noting that the super-resolution results are, however, not only limited to human consumption but also to industrial applications, such as control and automation [53, 54], object detection, and feature extraction. With the proposed method generating promising objective results, we hope that it may as well suit these other disciplines.

5 Conclusions

In this paper, we have presented an adaptive multi-frame super-resolution model that sufficiently restores fine image details. The new method incorporates a spatially varying regularizing term that updates its value according to the local image features—linear isotropic in flat regions and nonlinear anisotropic near edges. This flexibility and adaptability makes the model generate promising results, objectively and quantitatively. Also, the proposed adaptive term includes a convolution operation with the Gaussian filter, and this allows the model to robustly emphasize critical and meaningful features. Experimental results visually demonstrate the strength of the new method that it reveals more information in the reconstructed images compared with other methods. Objectively, we have shown that the method generates promising values of the quality metrics (PSNR, ESIM, and MSSIM).



Fig. 9 Simulated super-resolution results of different methods

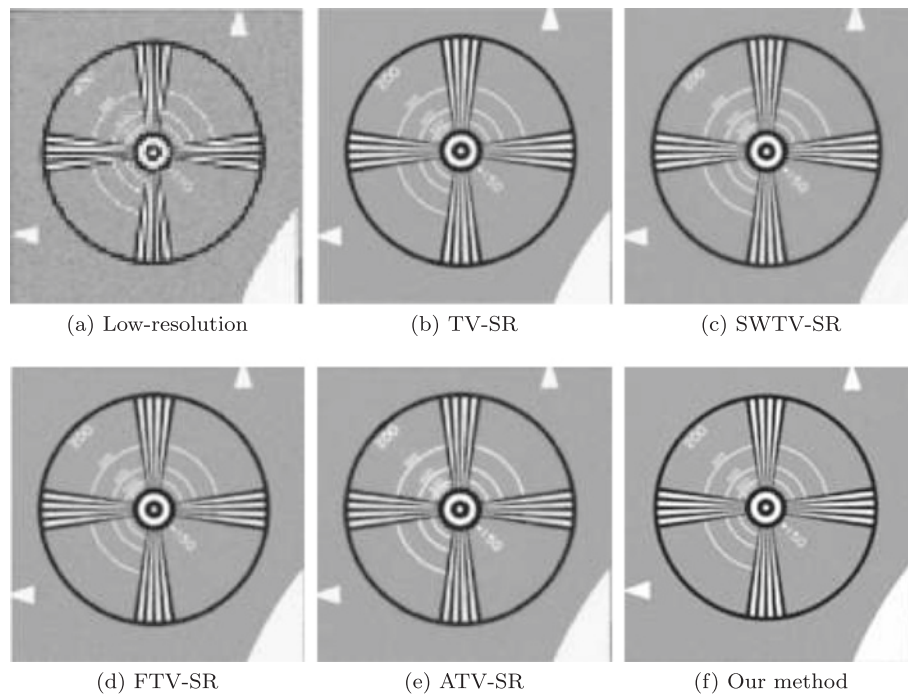


Fig. 10 Super-resolution results of different methods for real video sequences: **a** low-resolution, **b** TV-SR, **c** SWTV-SR, **d** FTV-SR, **e** ATV-SR, and **f** our method

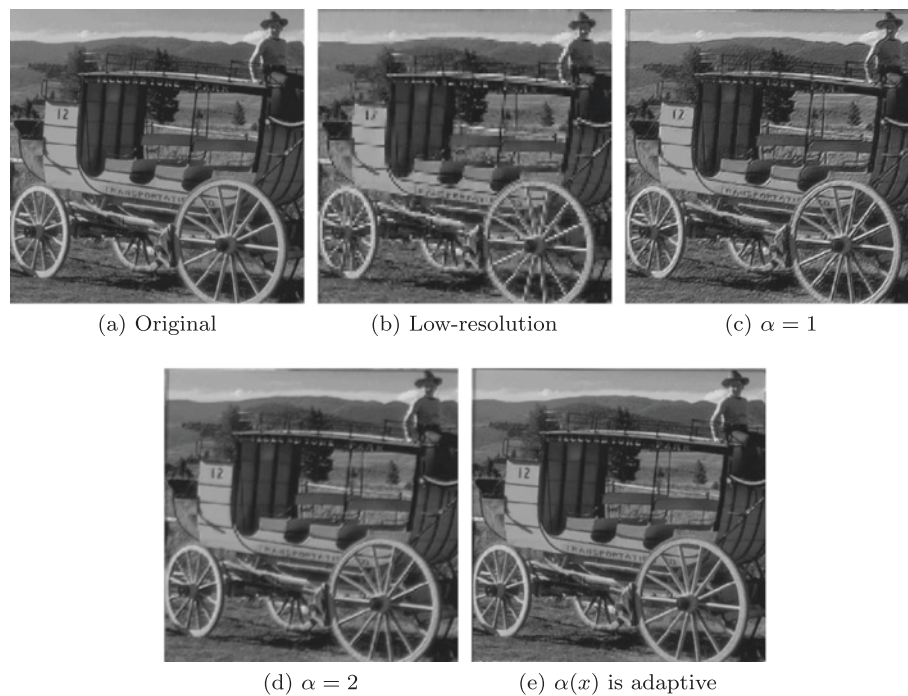


Fig. 11 Super-resolution results of our method for various conditions of the variable exponent, $\alpha(x)$: **a** original, **b** low-resolution, **c** $\alpha = 1$, **d** $\alpha = 2$, and **e** $\alpha(x)$

Table 1 Peak signal-to-noise ratio (PSNR) measurements of different super-resolution methods

Image	TV-SR	SWTV-SR	FTV-SR	ATV-SR	Our method
Boat	27.73	28.11	28.87	28.05	28.96
Bridge	28.02	29.33	30.15	30.21	31.29
Building	27.43	28.20	27.98	28.74	29.00
Fish	29.12	30.77	30.81	31.00	30.51
Goldhill	28.56	28.90	29.20	30.13	30.78
Lena	27.01	28.21	28.21	29.65	30.71
Mandrill	28.29	29.61	30.33	30.36	30.87
Wheel	29.98	31.05	30.89	31.67	31.59
Mean	28.27	29.27	29.56	29.98	30.46

Table 2 Edge similarity (ESIM) measurements of different super-resolution methods

Image	TV-SR	SWTV-SR	FTV-SR	ATV-SR	Our method
Boat	12.03	12.65	13.09	13.56	13.97
Bridge	13.17	14.00	13.88	14.23	14.78
Building	11.61	12.20	12.09	13.90	13.54
Fish	12.32	13.10	13.71	14.67	14.98
Goldhill	13.54	14.00	13.81	14.01	13.97
Lena	12.70	13.08	13.81	13.56	14.23
Mandrill	13.10	13.45	13.04	14.01	14.53
Wheel	14.01	14.21	13.89	13.22	14.23
Mean	12.81	13.34	13.42	13.90	14.28

Table 3 Mean structural similarity (MSSIM) measurements of different super-resolution methods

Image	TV-SR	SWTV-SR	FTV-SR	ATV-SR	Our method
Boat	0.7812	0.7709	0.7890	0.8102	0.8091
Bridge	0.7901	0.8093	0.8312	0.8200	0.8712
Building	0.7002	0.7856	0.7201	0.7412	0.7554
Fish	0.7812	0.8231	0.8660	0.8700	0.8798
Goldhill	0.7992	0.7085	0.7233	0.7310	0.7620
Lena	0.8441	0.8552	0.8301	0.8400	0.8577
Mandrill	0.8333	0.7898	0.7996	0.8206	0.8790
Wheel	0.7802	0.7790	0.7980	0.7912	0.7993
Mean	0.7887	0.7902	0.7947	0.8030	0.8227

Table 4 Algorithmic CPU times (in seconds) of different super-resolution methods

Image	TV-SR	SWTV-SR	FTV-SR	ATV-SR	Our method
Boat	7.23	8.21	10.19	5.11	6.54
Bridge	8.01	7.40	8.56	4.23	6.35
Building	7.99	7.48	8.32	6.90	5.02
Fish	7.31	6.55	9.02	5.34	5.91
Goldhill	7.87	7.32	8.67	5.40	6.78
Lena	5.12	8.90	10.00	7.98	5.07
Mandrill	9.05	8.92	7.35	4.77	4.33
Wheel	8.69	6.00	7.41	5.93	6.32
Mean	7.66	7.60	8.69	5.71	5.79

In the future, we are contemplating the possibilities of extending our method to other fields, such as in medical imaging. For example, doctors in sonography require high-quality ultrasound images to provide accurate treatments to patients. The current instruments produce low-quality images that are heavily degraded by multiplicative noise. The problem can be approached in a variety of ways. In the context of the new method, three important processes that may help to address the problem are (1) modifying the prior to cover multiplicative noise, (2) transforming the model into the three-dimensional space to comprehensively treat the ultrasound images, and (3) implementing the model using more accurate and fast numerical schemes that support real-time parallel computing.

Competing interests

The authors declare that they have no competing interests.

Author details

¹Department of Electronics and Telecommunication Engineering, College of Information and Communication Technologies, University of Dar es Salaam, P. O. Box 35194, 00255 Dar es Salaam, Tanzania. ²Research Institute of Intelligent Control and Systems, Harbin Institute of Technology, 150001 Harbin, China. ³Department of Mathematics, Egerton University, P. O. Box 536, 00254 Egerton, Nakuru, Kenya.

Received: 23 October 2014 Accepted: 24 June 2015

Published online: 28 July 2015

References

- B Goldlücke, M Aubry, K Kolev, D Cremers, A super-resolution framework for high-accuracy multiview reconstruction. *Int. J. Comput. Vis.* **106**(2), 172–191 (2014)
- S Park, M Park, M Kang, Super-resolution image reconstruction: a technical overview. *Signal Process. Mag. IEEE*. **20**(3), 21–36 (2003)
- B Maiseli, C Wu, J Mei, Q Liu, H Gao, A robust super-resolution method with improved high-frequency components estimation and aliasing correction capabilities. *J. Franklin Institute*. **351**, 513–527 (2014)
- RS Babu, KS Murthy, A survey on the methods of super-resolution image reconstruction. *Int. J. Comput. Appl.* **15**(2), 1–6 (2011)
- K Nasrollahi, TB Moeslund, Super-resolution: a comprehensive survey. *Mach. Vis. Appl.* **25**, 1423–1468 (2014)
- Y Zhou, Z Tang, X Hu, Fast single image super resolution reconstruction via image separation. *J. Netw.* **9**(7), 1811–1818 (2014)
- Y Tang, P Yan, Y Yuan, X Li, Single-image super-resolution via local learning. *Int. J. Mach. Learn. Cybernet.* **2**, 15–23 (2011)
- W Wu, Z Liu, X He, W Gueaieb, Single-image super-resolution based on Markov random field and contourlet transform. *J. Electronic Imaging*. **20**(2), 023005–023005 (2011)
- KI Kim, Y Kwon, Single-image super-resolution using sparse regression and natural image prior. *Pattern Anal. Mach. Intell. IEEE Trans.* **32**(6), 1127–1133 (2010)
- WT Freeman, TR Jones, EC Pasztor, Example-based super-resolution. *Comput. Graphics Appl. IEEE*. **22**(2), 56–65 (2002)
- K Zhang, X Gao, X Li, D Tao, Partially supervised neighbor embedding for example-based image super-resolution. *Selected Topics Signal Process IEEE J.* **5**(2), 230–239 (2011)
- C Kim, K Choi, JB Ra, Example-based super-resolution via structure analysis of patches. *Signal Process. Lett. IEEE*. **20**(4), 407–410 (2013)
- Z Xiong, D Xu, X Sun, F Wu, Example-based super-resolution with soft information and decision. *Multimedia IEEE Trans.* **15**(6), 1458–1465 (2013)
- CY Yang, JB Huang, MH Yang, in *Computer Vision—ACCV 2010*. Exploiting self-similarities for single frame super-resolution (Springer Queenstown, New Zealand, 2011), pp. 497–510
- O Mac Aodha, ND Campbell, A Nair, GJ Brostow, in *Computer Vision—ECCV 2012*. Patch based synthesis for single depth image super-resolution (Springer Florence, Italy, 2012), pp. 71–84
- X Gao, K Zhang, D Tao, X Li, Joint learning for single-image super-resolution via a coupled constraint. *Image Process IEEE Trans.* **21**(2), 469–480 (2012)
- J Li, X Peng, in *Information Science and Technology (ICIST) 2012 International Conference on*. Single-frame image super-resolution through gradient learning (IEEE Hubei, 2012), pp. 810–815
- K Zhang, X Gao, D Tao, X Li, Single image super-resolution with non-local means and steering kernel regression. *Image Process IEEE Trans.* **21**(11), 4544–4556 (2012)
- M Yang, Y Wang, *A self-learning approach to single image super-resolution*, vol. 15. (IEEE Transactions on Multimedia, 2013), pp. 498–508
- K Zhang, X Gao, D Tao, X Li, Single image super-resolution with multiscale similarity learning. *Neural Netw. Learn. Syst. IEEE Trans.* **24**(10), 1648–1659 (2013)
- X Li, Y Hu, X Gao, D Tao, B Ning, A multi-frame image super-resolution method. *Signal Processing*. **90**(2), 405–414 (2010)
- BJ Maiseli, Q Liu, OA Elisha, H Gao, Adaptive Charbonnier superresolution method with robust edge preservation capabilities. *J. Electronic Imaging*. **22**(4), 043027–043027 (2013)
- B Maiseli, O Elisha, J Mei, H Gao, *Edge preservation image enlargement and enhancement method based on the adaptive Perona–Malik non-linear diffusion model*, vol. 8. (IET Image Processing, 2014), pp. 753–760
- XD Zhao, ZF Zhou, JZ Cao, L Ren, GS Liu, H Wang, DS Wu, JH Tan, Multi-frame super-resolution reconstruction algorithm based on diffusion tensor regularization term. *Appl. Mech. Mater.* **543**, 2828–2832 (2014)
- M Elad, A Feuer, Restoration of a single superresolution image from several blurred, noisy, and undersampled measured images. *Image Process IEEE Trans.* **6**(12), 1646–1658 (1997)
- N Nguyen, P Milanfar, G Golub, A computationally efficient superresolution image reconstruction algorithm. *Image Process IEEE Trans.* **10**(4), 573–583 (2001)
- D Rajan, S Chaudhuri, An MRF-based approach to generation of super-resolution images from blurred observations. *J. Math. Imaging Vis.* **16**, 5–15 (2002)
- A Kanemura, Si Maeda, S Ishii, Superresolution with compound Markov random fields via the variational EM algorithm. *Neural Netw.* **22**(7), 1025–1034 (2009)
- KV Suresh, GM Kumar, A Rajagopalan, Superresolution of license plates in real traffic videos. *Intell. Transportation Syst. IEEE Trans.* **8**(2), 321–331 (2007)
- W Zeng, X Lu, A generalized DAMRF image modeling for superresolution of license plates. *Intell. Transport. Syst. IEEE Trans.* **13**(2), 828–837 (2012)
- S Mallat, G Yu, Super-resolution with sparse mixing estimators. *Image Process IEEE Trans.* **19**(11), 2889–2900 (2010)
- W Dong, D Zhang, G Shi, X Wu, Image deblurring and super-resolution by adaptive sparse domain selection and adaptive regularization. *Image Process IEEE Trans.* **20**(7), 1838–1857 (2011)
- A Marquina, SJ Osher, Image super-resolution by TV-regularization and Bregman iteration. *J. Scientific Comput.* **37**(3), 367–382 (2008)
- MK Ng, H Shen, EY Lam, L Zhang, A total variation regularization based super-resolution reconstruction algorithm for digital video. *EURASIP J. Adv. Signal Process.* **2007** (2007). doi:10.1155/2007/74585
- S Farsiu, MD Robinson, M Elad, P Milanfar, Fast and robust multiframe super resolution. *Image Process IEEE Trans.* **13**(10), 1327–1344 (2004)
- LI Rudin, S Osher, E Fatemi, Nonlinear total variation based noise removal algorithms. *Physica D: Nonlinear Phenomena*. **60**, 259–268 (1992)
- F Knoll, K Bredies, T Pock, R Stollberger, Second order total generalized variation (TGV) for MRI. *Magnet Resonance Med.* **65**(2), 480–491 (2011)
- P Getreuer, Total variation inpainting using split Bregman. *Image Process Line.* **2**, 147–157 (2012)
- A Bini, M Bhat, A nonlinear level set model for image deblurring and denoising. *Visual Comput.* **30**(3), 311–325 (2014)
- W Zeng, X Lu, S Fei, *Image super-resolution employing a spatial adaptive prior model*, vol. 162, (2015), pp. 218–233
- Z Ren, C He, Q Zhang, Fractional order total variation regularization for image super-resolution. *Signal Process.* **93**(9), 2408–2421 (2013)
- Q Yuan, L Zhang, H Shen, Multiframe super-resolution employing a spatially weighted total variation model. *Circ. Syst. Video Technol. IEEE Trans.* **22**(3), 379–392 (2012)
- EA Ogada, Z Guo, B Wu, in *Abstract and Applied Analysis, Volume 2014*. An alternative variational framework for image denoising (Hindawi

Publishing Corporation 410 Park Avenue 15th Floor, #287 pmb New York, NY 10022 USA, 2014)

44. P Perona, J Malik, Scale-space and edge detection using anisotropic diffusion. *Pattern Anal. Mach. Intell. IEEE Trans.* **12**(7), 629–639 (1990)
45. Z Guo, J Sun, D Zhang, B Wu, Adaptive Perona–Malik model based on the variable exponent for image denoising. *Image Process IEEE Trans.* **21**(3), 958–967 (2012)
46. S Levine, Y Chen, J Stanich, Image restoration via nonstandard diffusion. Duquesne University, Department of Mathematics and Computer Science Technical Report :04–01 (2004)
47. J Weickert, *Anisotropic diffusion in image processing, Volume 1*. (Teubner Stuttgart, 1998)
48. R Courant, K Friedrichs, H Lewy, On the partial difference equations of mathematical physics. *IBM J. Res. Dev.* **11**(2), 215–234 (1967)
49. Z Wang, AC Bovik, Mean squared error: love it or leave it? A new look at signal fidelity measures. *Signal Process Mag. IEEE.* **26**, 98–117 (2009)
50. A Tanchenko, Visual-PSNR measure of image quality. *J. Visual Commun. Image Representation.* **25**(5), 874–878 (2014)
51. Z Wang, AC Bovik, HR Sheikh, EP Simoncelli, Image quality assessment: from error visibility to structural similarity. *Image Process IEEE Trans.* **13**(4), 600–612 (2004)
52. S Durand, J Fadili, M Nikolova, Multiplicative noise removal using L1 fidelity on frame coefficients. *J. Math. Imaging Vis.* **36**(3), 201–226 (2010)
53. S Yin, X Li, H Gao, O Kaynak, *Data-based techniques focused on modern industry: an overview*, vol. 62. (IEEE Transactions on Industrial Electronics, 2015), pp. 657–667
54. S Yin, SX Ding, X Xie, H Luo, A review on basic data-driven approaches for industrial process monitoring. *IEEE Trans. Industrial Electronics.* **61**(11), 6418–6428 (2014)

Submit your manuscript to a SpringerOpen[®] journal and benefit from:

- Convenient online submission
- Rigorous peer review
- Immediate publication on acceptance
- Open access: articles freely available online
- High visibility within the field
- Retaining the copyright to your article

Submit your next manuscript at ► springeropen.com