

CHAPTER 9

Inverse Problems

September 9, 2026

Inverse problems seek to reconstruct a signal from measurements that blur, omit, or mix its values. Direct inversion can amplify noise or leave several signals indistinguishable, so reconstruction must incorporate assumptions about the unknown signal. We study how spectral and variational regularization stabilize inversion, derive error bounds for quadratic penalties, and develop methods for deconvolution, inpainting, and tomography.

The main references for this chapter are [2, 3, 1].

9.1 Regularization of Inverse Problems

We extend the convex regularization introduced in Chapter 8 to linear measurement operators.

We consider a bounded linear map $\Phi : \mathcal{S} \rightarrow \mathcal{H}$, where the signal space $\mathcal{S}$ is a Hilbert space or, more generally, a Banach space. The data space $\mathcal{H}$ is a Hilbert space. The operator models acquisition: an unknown high-resolution signal $f_0 \in \mathcal{S}$ gives rise to the noisy observation

$$y = \Phi f_0 + w \in \mathcal{H}$$

where $w \in \mathcal{H}$ represents acquisition noise. Here the noise is deterministic; we assume only that $\|w\|_{\mathcal{H}}$ is bounded.

An acquisition device records finitely many observations, so applications usually take $\mathcal{H} = \mathbb{R}^P$. The number P may be small relative to the desired reconstruction dimension. For numerical implementation, we discretize the signal space as $\mathcal{S} = \mathbb{R}^N$, where N is the number of grid points. Usually the grid is much finer than the measurement set, so $N \gg P$. Infinite-dimensional function spaces remain useful for modeling the unknown f_0 and analyzing recovery, particularly when choosing the signal space $\mathcal{S}$.

Direct inversion may be impossible because Φ is not injective or the data do not lie in its range. Even when an inverse exists on the range, it may have a large norm or be unbounded. Applying it to noisy data amplifies the perturbation: formally, $\Phi^{-1}y = f_0 + \Phi^{-1}w$.

We now give a few representative examples of forward operators Φ .

Denoising. The identity operator $\Phi = \text{Id}_{\mathcal{S}}$ with $\mathcal{S} = \mathcal{H}$ gives the denoising problem studied in Chapters 7 and 8.

De-blurring and super-resolution. For a general operator Φ , recovering f_0 requires both inversion and denoising. These goals often conflict because inversion amplifies noise. For example, in deblurring, Φ is a translation-invariant operator corresponding to low-pass filtering with a kernel h

$$\Phi f = f \star h. \tag{9.1}$$

A periodic model takes $\mathcal{S} = \mathcal{H} = L^2(\mathbb{T}^d)$; see Proposition 2.3. In practice, the blurred signal is sampled on a grid: $\Phi f = \{(f \star h)(x_k) ; 0 \leq k < P\}$. Figure 9.1, middle, shows the resulting low-resolution image Φf_0 . Inverting such operators is useful for increasing the resolution of digital photographs and videos.

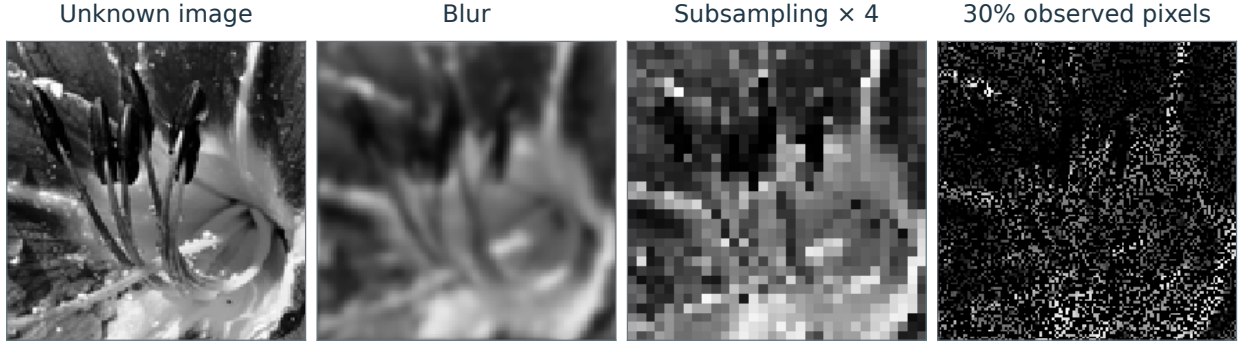


Figure 9.1. Examples of forward operators in inverse problems.

Interpolation and inpainting. Inpainting reconstructs missing pixels. We model the masking operation by an operator that is diagonal in the spatial domain:

$$(\Phi f)(x) = \begin{cases} 0 & \text{if } x \in \Omega, \\ f(x) & \text{if } x \notin \Omega. \end{cases} \quad (9.2)$$

Here Ω is the missing region, either a subset of $[0, 1]^d$ in the continuous model or a set of pixel indices in the discrete model. Figure 9.1, right, shows an example of a damaged image Φf_0 .

Medical imaging. Many medical imaging devices provide indirect access to the signal of interest, with acquisition processes that can be approximated by a linear operator Φ . For a tomographic scanner, the Radon transform models acquisition. The Fourier slice theorem relates these measurements to Fourier data along radial lines. An idealized magnetic resonance imaging (MRI) model also uses partial Fourier measurements:

$$\Phi f = \left\{ \hat{f}(x) ; x \in \Omega \right\}. \quad (9.3)$$

Here Ω consists of radial lines for tomography or acquisition trajectories such as spirals for MRI.

Electroencephalography (EEG) and magnetoencephalography (MEG) also infer internal activity from measurements at sensors outside the sources. Abstractly, such a linearized model can be written as $(\Phi f)(s) = \int K(s, x) f(x) dx$, where the kernel depends on the geometry and physical model. Unlike a simple image blur, this kernel need not be translation invariant.

Regression for supervised learning. Linear supervised learning is closely related to the imaging problems studied here, although it uses different notation. Section 13.3 develops the connection between regression and inverse problems. In statistical learning, the data are n pairs $(x_i, y_i)_{i=1}^n$ with feature vectors $x_i \in \mathbb{R}^p$. A linear prediction model has the form $y_i = \langle \beta, x_i \rangle$, with unknown parameter $\beta \in \mathbb{R}^p$. Placing the vectors x_i in the rows of $X \in \mathbb{R}^{n \times p}$ gives the approximate system $X\beta \approx y$. This corresponds to the inverse problem $\Phi f = y$ under the substitutions $\Phi \mapsto X$ and $f \mapsto \beta$, with dimensions $(P, N) \rightarrow (n, p)$. In statistical learning, the model need not be correctly specified, and the design matrix X is often random. Its sampling fluctuations must then be controlled as $n \rightarrow \infty$.

Ridge regression estimates the parameter by solving the normalized problem

$$\min_{\beta} \frac{1}{2n} \|X\beta - y\|^2 + \lambda \|\beta\|^2$$

The corresponding empirical quantities are therefore $\frac{1}{n} X^* X \sim \Phi^* \Phi$ for the covariance and $\frac{1}{n} X^* y \sim \Phi^* y$ for the data term. For independent observations in fixed dimension and with suitable finite moments, these empirical quantities fluctuate around their expectations at order $n^{-1/2}$. In growing dimension, the relevant norm bounds also depend on the dimension and tail assumptions.

9.2 Warmup: Oracle Linear Inversion

We first seek an oracle estimator among inversion methods that act diagonally in a fixed basis, allowing the gains to depend on the unknown clean signal. The optimization then separates into one scalar problem per coefficient. Later we will obtain such estimators from a single optimization problem on the signal space.

Assume that the forward operator is diagonal in an orthonormal basis (ψ_k) :

$$\Phi f = \sum_k \varphi_k \langle f, \psi_k \rangle \psi_k,$$

where the coefficients φ_k typically represent attenuation. The observation model is

$$Y = \Phi f_0 + w$$

where w is Gaussian white noise with variance σ^2 . In a real orthonormal basis, its coefficients are independent centered Gaussians. In a complex Fourier basis they still have mean zero and second moment σ^2 , which suffices for the risk calculation below. We work in finite dimension; infinite-dimensional white noise requires a generalized random-process interpretation. For instance, if $(\psi_k)_k$ is the Fourier basis, this corresponds to a deconvolution problem $\Phi f = h \star f$ where $\hat{h}_k = \varphi_k$. We consider the same diagonal estimator $\tilde{f}$ as in the denoising case:

$$\tilde{f} := \sum_k \lambda_k \langle Y, \psi_k \rangle \psi_k.$$

We minimize the mean squared error with respect to λ :

$$\mathbb{E}_w \|\tilde{f} - f_0\|^2 = \sum_k \mathbb{E} |\langle \tilde{f} - f_0, \psi_k \rangle|^2 = \sum_k \mathbb{E} |\lambda_k (\varphi_k c_k + \langle w, \psi_k \rangle) - c_k|^2.$$

Writing $c_k = \langle f_0, \psi_k \rangle$ and using the zero mean and variance of the noise,

$$\mathbb{E}_w \|\tilde{f} - f_0\|^2 = \sum_k |c_k|^2 |\varphi_k \lambda_k - 1|^2 + |\lambda_k|^2 \sigma^2.$$

The optimal coefficient of the oracle estimator solves

$$\min_{\lambda_k} |c_k|^2 |\varphi_k \lambda_k - 1|^2 + |\lambda_k|^2 \sigma^2,$$

i.e., it satisfies

$$|c_k|^2 \bar{\varphi}_k (\varphi_k \lambda_k - 1) + \sigma^2 \lambda_k = 0,$$

which gives

$$\lambda_k = \frac{|c_k|^2 \bar{\varphi}_k}{|c_k|^2 |\varphi_k|^2 + \sigma^2}.$$

For $\varphi_k \neq 0$ and $c_k \neq 0$, the oracle gains approach direct inversion as the noise vanishes:

$$\lambda_k \longrightarrow \frac{1}{\varphi_k} \quad \text{as } \sigma \rightarrow 0.$$

As σ increases, inversion is attenuated. For a simple Fourier-domain example, let $k \geq 1$, $|c_k|^2 = k^{-\alpha}$ and $\varphi_k = k^{-\beta/2}$, with $\alpha, \beta > 0$. Then

$$\lambda_k = \frac{k^{\beta/2}}{1 + \sigma^2 k^{\alpha+\beta}}.$$

At low frequencies this filter approximately compensates for attenuation, whereas at high frequencies it suppresses noise. Treating k as a positive real variable, its maximum occurs at

$$k^* = \left(\frac{\beta}{(2\alpha + \beta)\sigma^2} \right)^{1/(\alpha+\beta)}.$$

On the discrete range $k \geq 1$, the maximum is attained near this value, or at the lowest frequency if $k^* < 1$.

9.3 Theoretical Study of Quadratic Regularization

We outline a standard approach to recovery guarantees in Hilbert spaces. The emphasis is on modeling assumptions expressed through source conditions, and on limitations of the squared-norm penalty, notably saturation of convergence rates and loss of components in $\ker(\Phi)$.

9.3.1 Singular Value Decomposition

Finite dimension. We begin with the finite-dimensional case $\Phi \in \mathbb{R}^{P \times N}$ so that $\mathcal{S} = \mathbb{R}^N$ and $\mathcal{H} = \mathbb{R}^P$ are Hilbert spaces. In this case, the singular value decomposition (SVD) provides a precise description of the operator and of linear inversion.

Proposition 9.1 (SVD). *There exist matrices $(U, V) \in \mathbb{R}^{P \times R} \times \mathbb{R}^{N \times R}$, with $R = \text{rank}(\Phi) = \dim(\text{Im}(\Phi))$, satisfying $U^\top U = V^\top V = \text{Id}_R$. Their orthonormal columns $(u_m)_{m=1}^R \subset \mathbb{R}^P$, $(v_m)_{m=1}^R \subset \mathbb{R}^N$ and singular values $(\sigma_m)_{m=1}^R$, with $\sigma_m > 0$, satisfy*

$$\Phi = U \text{diag}_m(\sigma_m) V^\top = \sum_{m=1}^R \sigma_m u_m v_m^\top. \quad (9.4)$$

Proof. To motivate the construction, suppose that $\Phi = U \Sigma V^\top$ with $\Sigma = \text{diag}_m(\sigma_m)$. Then $\Phi \Phi^\top = U \Sigma^2 U^\top$, which determines the left singular vectors, and $V^\top = \Sigma^{-1} U^\top \Phi$ determines the right ones. Since $\Phi \Phi^\top$ is symmetric and positive semidefinite, its eigendecomposition on its range is $\Phi \Phi^\top = U \Sigma^2 U^\top$, where $\Sigma = \text{diag}_{m=1}^R(\sigma_m)$ with $\sigma_m > 0$. Define $V \stackrel{\text{def}}{=} \Phi^\top U \Sigma^{-1}$. One then verifies that

$$V^\top V = (\Sigma^{-1} U^\top \Phi)(\Phi^\top U \Sigma^{-1}) = \Sigma^{-1} U^\top (U \Sigma^2 U^\top) U \Sigma^{-1} = \text{Id}_R \quad \text{and} \quad U \Sigma V^\top = U \Sigma \Sigma^{-1} U^\top \Phi = \Phi.$$

□

The theorem also holds for complex matrices, with $^\top$ replaced by * . Expression (9.4) describes Φ as a sum of rank-1 matrices $u_m v_m^\top$. One usually orders the singular values $(\sigma_m)_m$ in decreasing order $\sigma_1 \geq \dots \geq \sigma_R$. If the positive singular values are distinct, the real reduced SVD is unique up to simultaneous sign changes of each pair (u_m, v_m) . In the complex case, simultaneous multiplication by a unit-modulus scalar replaces the sign change.

The columns of U form an orthonormal basis of $\text{Im}(\Phi)$, while the columns of V form an orthonormal basis of $\text{Im}(\Phi^\top) = \ker(\Phi)^\perp$. The decomposition (9.4) is called the reduced SVD because it retains only the R nonzero singular values. The full SVD completes the columns of U and V to orthonormal bases of $\mathbb{R}^P$ and $\mathbb{R}^N$, respectively. The diagonal factor Σ then has size $P \times N$.

For a periodic discrete convolution $\Phi f = h \star f$, let F be the unitary discrete Fourier matrix. Then

$$\Phi = F^* \text{diag}(\hat{h}_m) F, \quad (9.5)$$

where $\hat{h}_m$ denotes the convolution multiplier (the unnormalized DFT of h for the usual discrete convolution). If $q_m = F^* e_m$ and $\hat{h}_m \neq 0$, one may take $v_m = q_m$, $u_m = (\hat{h}_m / |\hat{h}_m|) q_m$, and $\sigma_m = |\hat{h}_m|$.

Computing the SVD of a dense matrix $\Phi \in \mathbb{R}^{N \times N}$ typically costs $O(N^3)$ operations.

Compact operators. The SVD extends to compact operators $\Phi : \mathcal{S} \rightarrow \mathcal{H}$ between separable Hilbert spaces. Compactness means that ΦB_1 is relatively compact, where $B_1 = \{s \in \mathcal{S} ; \|s\| \leq 1\}$ is the unit ball. Equivalently, every sequence $(\Phi s_k)_k$ with $s_k \in B_1$ has a convergent subsequence. In infinite dimension, the identity operator $\Phi : \mathcal{S} \rightarrow \mathcal{S}$ is not compact.

Equivalently, compact operators Φ admit an expansion analogous to (9.4):

$$\Phi = \sum_{m=1}^{+\infty} \sigma_m \langle \cdot, v_m \rangle u_m \quad (9.6)$$

Here (u_m) and (v_m) are orthonormal systems in $\mathcal{H}$ and $\mathcal{S}$, respectively, and the positive singular values decrease to zero; finite-rank operators give a finite sum. Convergence in (9.6) holds in the operator norm induced by the norms on $\mathcal{S}$ and $\mathcal{H}$

$$\|\Phi\|_{\mathcal{L}(\mathcal{S}, \mathcal{H})} \stackrel{\text{def}}{=} \sup_{\|u\|_{\mathcal{S}} \leq 1} \|\Phi u\|_{\mathcal{H}}.$$

For a nonzero operator Φ with the expansion (9.6), $\|\Phi\|_{\mathcal{L}(\mathcal{S}, \mathcal{H})} = \sigma_1$.

If only R singular values are positive, then Φ has finite rank $R = \dim(\text{Im}(\Phi))$. Squared-norm regularization restricts the reconstruction to $\ker(\Phi)^\perp$, a space of dimension R , so the problem reduces to finite dimension. Nonlinear methods can also recover components in $\ker(\Phi)$ under suitable priors; this is sometimes called super-resolution. Between Hilbert spaces, compact operators are exactly the operator-norm limits of finite-rank operators. This approximation property need not hold for arbitrary Banach spaces, although the definition of compactness itself extends to them.

Examples include integral operators on $L^2(\Omega)$ with a continuous kernel $k(x, y)$ defined for $(x, y) \in \Omega \times \Omega$, where Ω is a compact subset of $\mathbb{R}^d$ or the torus $\mathbb{T}^d$:

$$(\Phi f)(x) = \int_{\Omega} k(x, y) f(y) dy$$

where dy is the Lebesgue measure. Convolution $\Phi f = f \star h$ on $\mathbb{T}^d = (\mathbb{R}/2\pi\mathbb{Z})^d$ generalizes (9.5). Its kernel $k(x, y) = h(x - y)$ is translation invariant, and the Fourier functions $q_m(x) = (2\pi)^{-d/2} e^{im \cdot x}$ diagonalize the operator. Its singular values are $\sigma_m = |\hat{h}_m|$, where $\hat{h}_m = \int_{\mathbb{T}^d} h(x) e^{-im \cdot x} dx$ denotes the convolution multiplier. Another example on $\Omega = [0, 1]$ is the integration operator $(\Phi f)(x) = \int_0^x f(y) dy$, which has the square-integrable kernel $k(x, y) = 1_{y \leq x}$ and is also compact.

Pseudoinverse. We first work in finite dimension, or more generally assume that $\text{Im}(\Phi)$ is closed. Even without noise, $\Phi f = y$ may be inconsistent, and a nontrivial kernel makes a solution nonunique. The Moore–Penrose pseudoinverse selects the least-squares solution of minimum norm

$$f^+ \stackrel{\text{def}}{=} \underset{\Phi f = y^+}{\text{argmin}} \|f\|_{\mathcal{S}} \quad \text{where} \quad y^+ = \text{Proj}_{\text{Im}(\Phi)}(y) = \underset{z \in \text{Im}(\Phi)}{\text{argmin}} \|y - z\|_{\mathcal{H}}.$$

The next proposition expresses the pseudoinverse through the SVD. Under injectivity or surjectivity, it also gives formulas involving linear systems with $\Phi\Phi^*$ or $\Phi^*\Phi$.

Proposition 9.2. *One has*

$$f^+ = \Phi^+ y \quad \text{where} \quad \Phi^+ = V \text{diag}_m(1/\sigma_m) U^*. \quad (9.7)$$

If $\text{Im}(\Phi) = \mathcal{H}$, one has $\Phi^+ = \Phi^*(\Phi\Phi^*)^{-1}$. If $\ker(\Phi) = \{0\}$, one has $\Phi^+ = (\Phi^*\Phi)^{-1}\Phi^*$.

Proof. The columns of U form an orthonormal basis of $\text{Im}(\Phi)$, so $y^+ = UU^*y$. Thus $\Phi f = y^+$ becomes $U\Sigma V^*f = UU^*y$, or $V^*f = \Sigma^{-1}U^*y$. In the orthogonal decomposition $f = f_0 + r$, with $f_0 \in \ker(\Phi)^\perp$ and $r \in \ker(\Phi)$, the data fix $f_0 = VV^*f = V\Sigma^{-1}U^*y = \Phi^+y$. Since $\|f\|^2 = \|f_0\|^2 + \|r\|^2$, minimizing the total norm amounts to minimizing $\|r\|$, which gives $r = 0$. If $\text{Im}(\Phi) = \mathcal{H}$, then $R = P$ so that $\Phi\Phi^* = U\Sigma^2U^*$ is the eigendecomposition of an invertible matrix, and $(\Phi\Phi^*)^{-1} = U\Sigma^{-2}U^*$. One then verifies $\Phi^*(\Phi\Phi^*)^{-1} = V\Sigma U^*U\Sigma^{-2}U^*$ which is the desired result. The second case follows similarly. $\square$

In infinite dimension, a compact operator of infinite rank has nonclosed range. Its pseudoinverse is unbounded and is defined only for data satisfying the corresponding square-summability condition on the inverse-weighted SVD coefficients. The projections and inverse formulas above cannot then be used on arbitrary data.

For convolution operators $\Phi f = f \star h$, the spectral formula gives

$$\Phi^+ y = y \star h^+ \quad \text{where} \quad \hat{h}_m^+ = \begin{cases} \hat{h}_m^{-1} & \text{if } \hat{h}_m \neq 0 \\ 0 & \text{if } \hat{h}_m = 0. \end{cases}$$

9.3.2 Tikhonov Regularization

Regularized inverse. For an ill-conditioned operator, applying the unregularized inverse to noisy data can be unstable, since

$$\Phi^+ y = \Phi^+ \Phi f_0 + \Phi^+ w = f_0^+ + \Phi^+ w \quad \text{where} \quad f_0^+ \stackrel{\text{def}}{=} \text{Proj}_{\ker(\Phi)^\perp}(f_0),$$

The recovery error is therefore $\|\Phi^+ y - f_0^+\| = \|\Phi^+ w\|$. It reaches $\|w\|/\sigma_R$ when $w \propto u_R$, so a small singular value produces a large amplification factor $1/\sigma_R$. In the infinite-rank case $R = +\infty$, the compact operator's pseudoinverse is unbounded. To control this amplification, replace Φ^+ by a regularized inverse of the form

$$\Phi_\lambda^+ = V \text{diag}_m(\mu_\lambda(\sigma_m)) U^* \quad (9.8)$$

where the spectral filter μ_λ depends on a parameter $\lambda > 0$ and satisfies the boundedness and consistency conditions

$$|\mu_\lambda(\sigma)| \leq C_\lambda < +\infty \quad \text{and} \quad \lim_{\lambda \rightarrow 0} \mu_\lambda(\sigma) = \frac{1}{\sigma}.$$

Figure 9.2, left, illustrates a capped reciprocal filter, which limits the amplification of small singular values.

Variational regularization. A typical regularized inverse is obtained by solving a penalized least-squares problem

$$f_\lambda \stackrel{\text{def.}}{=} \operatorname{argmin}_{f \in \mathcal{S}} \|y - \Phi f\|_{\mathcal{H}}^2 + \lambda J(f) \quad (9.9)$$

where J is a regularization functional. For general convex priors, lower semicontinuity, rather than continuity, is the natural assumption; existence also requires compactness or coercivity conditions. The simplest example is the quadratic norm $J = \|\cdot\|_{\mathcal{S}}^2$,

$$f_\lambda \stackrel{\text{def.}}{=} \operatorname{argmin}_{f \in \mathcal{S}} \|y - \Phi f\|_{\mathcal{H}}^2 + \lambda \|f\|_{\mathcal{S}}^2 \quad (9.10)$$

Proposition 9.3 shows that this estimator is a special case of (9.8). Its normal equations give

$$f_\lambda = (\Phi^* \Phi + \lambda \operatorname{Id}_{\mathcal{S}})^{-1} \Phi^* y. \quad (9.11)$$

This shows that $f_\lambda \in \operatorname{Im}(\Phi^*) \subset \ker(\Phi)^\perp$, and that it depends linearly on y .

Proposition 9.3. *The solution of (9.10) is $f_\lambda = \Phi_\lambda^+ y$, with the regularized inverse (9.8) defined by the filter*

$$\forall \sigma \geq 0, \quad \mu_\lambda(\sigma) = \frac{\sigma}{\sigma^2 + \lambda}.$$

Proof. The normal equations are $(\Phi^* \Phi + \lambda \operatorname{Id}_{\mathcal{S}}) f_\lambda = \Phi^* y$. Taking components along v_m gives

$$(\sigma_m^2 + \lambda) \langle f_\lambda, v_m \rangle = \sigma_m \langle y, u_m \rangle.$$

The component in $\ker(\Phi)$ vanishes because $\lambda > 0$. Thus

$$f_\lambda = \sum_m \frac{\sigma_m}{\sigma_m^2 + \lambda} \langle y, u_m \rangle v_m,$$

which proves the formula. □

For convolution $\Phi f = f \star h$, the FFT computes the regularized inverse in $O(N \log(N))$ operations:

$$\hat{f}_{\lambda,m} = \frac{\hat{h}_m^*}{|\hat{h}_m|^2 + \lambda} \hat{y}_m.$$

Figure 9.2 compares quadratic regularization (9.11) (right) with a capped reciprocal filter (left).

We choose λ according to the noise level and seek a rate for $f_\lambda \rightarrow f_0$. This requires $f_0 = f_0^+$, equivalently $f_0 \in \overline{\operatorname{Im}(\Phi^*)} = \ker(\Phi)^\perp$: squared-norm regularization always produces $f_\lambda \in \operatorname{Im}(\Phi^*) \subset \ker(\Phi)^\perp$, so it cannot recover a kernel component of f_0 . Nonlinear priors can recover such components when additional structure makes the signal identifiable.

Source condition. To obtain a convergence rate, we impose a source condition of order β , which reads

$$f_0 \in \operatorname{Im}((\Phi^* \Phi)^{\beta/2}) = \operatorname{Im}(V \operatorname{diag}(\sigma_m^\beta) V^*). \quad (9.12)$$

For a fixed bound on the source vector, larger β imposes stronger decay along directions with small singular values and makes stable inversion easier. The condition asserts the existence of $z \in \mathcal{S}$ such that $f_0 = V \operatorname{diag}(\sigma_m^\beta) V^* z$. Choose the source vector of minimum norm, $z = V \operatorname{diag}(\sigma_m^{-\beta}) V^* f_0$, and impose $\|z\| \leq \rho$ for some $\rho > 0$. This is equivalent to the coefficient bound

$$\sum_m \sigma_m^{-2\beta} |\langle f_0, v_m \rangle|^2 \leq \rho^2 < +\infty. \quad (S_{\beta,\rho})$$

The assumptions $\beta > 0$ and $f_0 \in \ker(\Phi)^\perp$ are part of the source condition; the coefficient bound alone leaves kernel components unconstrained. The bound resembles the Sobolev constraint in 7.6. For example, consider a low-pass convolution $\Phi f = f \star h$ where the kernel h has a polynomially decaying multiplier: $|\hat{h}_m| \sim 1/m^\alpha$ for large m . Since the vectors v_m are Fourier modes, $(S_{\beta,\rho})$ gives, up to constants and the Fourier normalization, a bound of the form

$$\sum_m \|m\|^{2\alpha\beta} |\hat{f}_m|^2 \leq \rho^2 < +\infty.$$

This is a Sobolev-type coefficient bound of smoothness order $\alpha\beta$.

Sublinear convergence speed. The following theorem gives the convergence rate implied by this source condition. With the noise bound $\|w\| \leq \delta$, recovery depends on the parameters (δ, ρ, β) . Assuming $f_0 \in \ker(\Phi)^\perp$, we study the convergence of f_λ to f_0 for data $y = \Phi f_0 + w$ as $\delta \rightarrow 0$. The analysis also determines how λ should depend on δ .

Theorem 9.4. *Let $\delta > 0$. Under the source condition $(S_{\beta,\rho})$ with $0 < \beta \leq 2$, the solution of (9.10) for $\|w\| \leq \delta$ satisfies*

$$\|f_\lambda - f_0\| \leq C \rho^{\frac{1}{\beta+1}} \delta^{\frac{\beta}{\beta+1}}$$

for a constant C which depends only on β , and for a choice

$$\lambda \sim \delta^{\frac{2}{\beta+1}} \rho^{-\frac{2}{\beta+1}}.$$

Proof. The source condition implies $f_0 \in \ker(\Phi)^\perp$. We decompose

$$f_\lambda = \Phi_\lambda^+(\Phi f_0 + w) = f_\lambda^0 + \Phi_\lambda^+ w \quad \text{where} \quad f_\lambda^0 \stackrel{\text{def}}{=} \Phi_\lambda^+(\Phi f_0),$$

Thus $f_\lambda = f_\lambda^0 + \Phi_\lambda^+ w$, and for any regularized inverse of the form (9.8)

$$\|f_\lambda - f_0\| \leq \|f_\lambda - f_\lambda^0\| + \|f_\lambda^0 - f_0\|. \quad (9.13)$$

The noise-propagation term $\|f_\lambda - f_\lambda^0\|$ measures the effect of measurement noise. For Tikhonov regularization it decreases as λ increases and is often called the variance term. The bias term $\|f_\lambda^0 - f_0\|$ is independent of the noise and measures the smoothing of f_0 . For Tikhonov regularization it increases with λ . Choosing λ therefore balances bias against noise amplification. Assuming

$$\forall \sigma \geq 0, \quad |\mu_\lambda(\sigma)| \leq C_\lambda$$

the variance term is bounded as

$$\|f_\lambda - f_\lambda^0\|^2 = \|\Phi_\lambda^+ w\|^2 = \sum_m \mu_\lambda(\sigma_m)^2 |w_m|^2 \leq C_\lambda^2 \|w\|_{\mathcal{H}}^2.$$

Using $\frac{|f_{0,m}|^2}{\sigma_m^{2\beta}} = |z_m|^2$, we bound the bias term as follows:

$$\|f_\lambda^0 - f_0\|^2 = \sum_m (1 - \mu_\lambda(\sigma_m)\sigma_m)^2 |f_{0,m}|^2 = \sum_m \left((1 - \mu_\lambda(\sigma_m)\sigma_m)\sigma_m^\beta \right)^2 \frac{|f_{0,m}|^2}{\sigma_m^{2\beta}} \leq D_{\lambda,\beta}^2 \rho^2 \quad (9.14)$$

where we assumed

$$\forall \sigma \geq 0, \quad |(1 - \mu_\lambda(\sigma)\sigma)\sigma^\beta| \leq D_{\lambda,\beta}. \quad (9.15)$$

For $\beta > 2$, the supremum over all $\sigma \geq 0$ is infinite. On the bounded spectrum of a fixed operator, the bias remains bounded, but the resulting rate saturates. Putting (9.14) and (9.15) together, one obtains

$$\|f_\lambda - f_0\| \leq C_\lambda \delta + D_{\lambda,\beta} \rho. \quad (9.16)$$

In the case of the regularization (9.10), one has $\mu_\lambda(\sigma) = \frac{\sigma}{\sigma^2 + \lambda}$, and thus $(1 - \mu_\lambda(\sigma)\sigma)\sigma^\beta = \frac{\lambda\sigma^\beta}{\sigma^2 + \lambda}$. For $\beta \leq 2$, one verifies (see Figure 9.2 and 9.3) that

$$C_\lambda = \frac{1}{2\sqrt{\lambda}} \quad \text{and} \quad D_{\lambda,\beta} = c_\beta \lambda^{\frac{\beta}{2}},$$

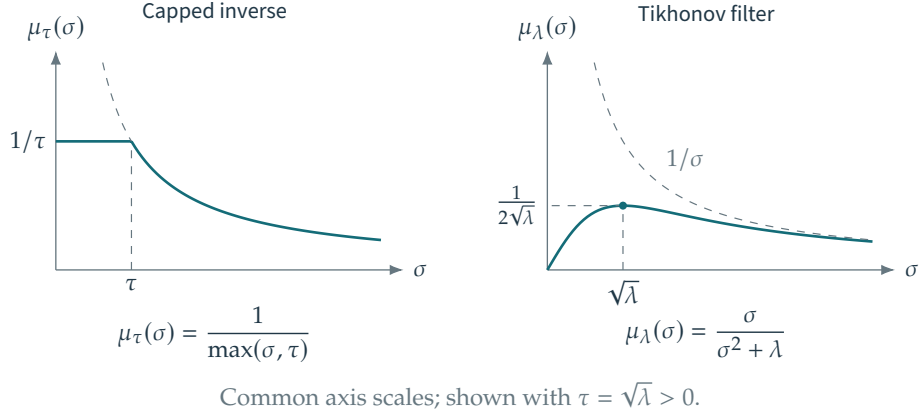


Figure 9.2. Left: capping reciprocal singular values. Right: the Tikhonov bound $\mu_\lambda(\sigma) \leq C_\lambda = \frac{1}{2\sqrt{\lambda}}$.

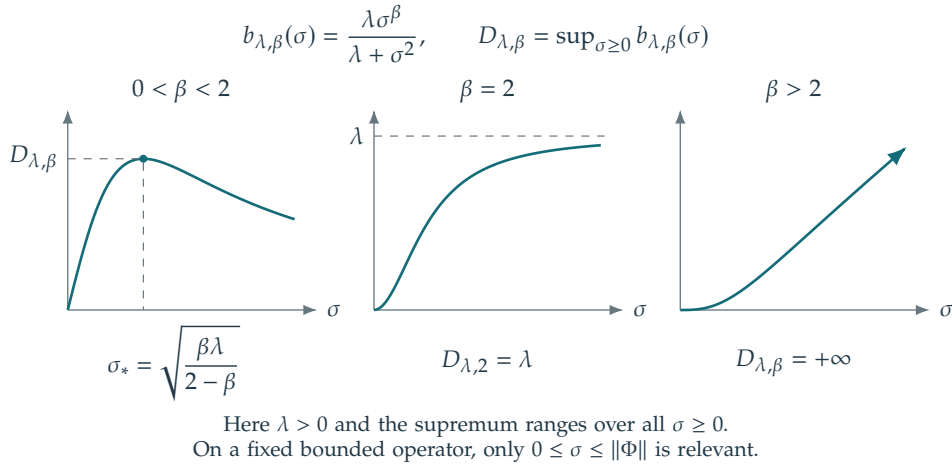


Figure 9.3. Bounding $\lambda \frac{\sigma^\beta}{\lambda + \sigma^2} \leq D_{\lambda,\beta}$.

for some constant c_β . Balancing the two terms in (9.16) gives the correct order; optimizing over λ would improve only the constant. Set $\frac{\delta}{\sqrt{\lambda}} = \lambda^{\frac{\beta}{2}} \rho$, which gives $\lambda = (\delta/\rho)^{\frac{2}{\beta+1}}$. Then

$$\|f_\lambda - f_0\| = O(\delta/\sqrt{\lambda}) = O(\delta(\delta/\rho)^{-\frac{1}{\beta+1}}) = O(\delta^{\frac{\beta}{\beta+1}} \rho^{\frac{1}{\beta+1}}).$$

□

Larger source orders $\beta \leq 2$ therefore give faster convergence as $\|w\|$ tends to zero. The rate in Theorem 9.4 saturates: choosing $\beta > 2$ gives the same general worst-case rate as $\beta = 2$. The best rate obtainable in this way is

$$\|f_\lambda - f_0\| = O(\rho^{\frac{1}{3}} \delta^{\frac{2}{3}}).$$

Alternative spectral filters μ_λ , together with sufficiently large β , can give the rate $\|f_\lambda - f_0\| = O(\delta^{1-\kappa})$ for arbitrarily small $\kappa > 0$. The capped inverse in Figure 9.2, left, avoids this finite-order saturation, whereas Tikhonov regularization is limited to source orders $\beta \leq 2$. Quadratic regularization is easier to implement: its variational formulation leads to a linear system and does not require an SVD. For compact operators with nonclosed range, no uniform linear rate follows over these source balls from finite-order source conditions alone. Linear rates are possible in finite dimension with a fixed smallest positive singular value, or under different structural assumptions, such as those used for sparse ℓ^1 regularization in Chapter 10.

9.4 Quadratic Regularization

We now turn from infinite-dimensional recovery theory to numerical methods in finite dimension.

Convex regularization. As in (9.9), we reconstruct a high-resolution image $f_0 \in \mathbb{R}^N$ from noisy measurements $y = \Phi f_0 + w \in \mathbb{R}^P$ by minimizing a convex regularized objective:

$$f_\lambda \in \operatorname{argmin}_{f \in \mathbb{R}^N} \mathcal{E}(f) \stackrel{\text{def.}}{=} \frac{1}{2} \|y - \Phi f\|^2 + \lambda J(f) \quad (9.17)$$

Here $\|y - \Phi f\|^2$ measures data fidelity, with $\|\cdot\|$ denoting the ℓ^2 norm on $\mathbb{R}^P$, and $J(f)$ is a convex penalty on $\mathbb{R}^N$.

The regularization parameter $\lambda > 0$ balances these two terms and can be difficult to choose in practice. In simulations with a known reference signal f_0 , the parameter can be calibrated by minimizing the reconstruction error $\|f_0 - \tilde{f}\|$. Such a reference is generally unavailable in applications.

For noiseless data, $w = 0$, we study small values of λ . As $\lambda \rightarrow 0$, minimizers of (9.17) approach solutions of a constrained problem under the assumptions below. We may assume $y \in \operatorname{Im}(\Phi)$ without changing the minimizers, since orthogonality gives

$$\|y - \Phi f\|^2 = \|y - \operatorname{Proj}_{\operatorname{Im}(\Phi)}(y)\|^2 + \|\operatorname{Proj}_{\operatorname{Im}(\Phi)}(y) - \Phi f\|^2$$

The first term is independent of f , so replacing y by $\operatorname{Proj}_{\operatorname{Im}(\Phi)}(y)$ in (9.17) changes only an additive constant.

Let us recall that a function J is coercive if

$$\lim_{\|f\| \rightarrow +\infty} J(f) = +\infty$$

i.e.

$$\forall K \in \mathbb{R}, \exists R > 0, \quad \|f\| \geq R \implies J(f) \geq K.$$

Equivalently, every sublevel set $\{f : J(f) \leq c\}$ is bounded. In finite dimension, lower semicontinuity of J then makes these sets compact.

Proposition 9.5. Assume that $J : \mathbb{R}^N \rightarrow \mathbb{R}$ is lower semicontinuous and coercive and that $y \in \operatorname{Im}(\Phi)$. For each λ , let f_λ solve (9.17). The family $(f_\lambda)_\lambda$ is bounded, and every accumulation point $f^\star$ as $\lambda \downarrow 0$ solves

$$f^\star \in \operatorname{argmin}_{f \in \mathbb{R}^N} \{J(f) : \Phi f = y\}. \quad (9.18)$$

Proof. Lower semicontinuity and coercivity ensure that J attains its minimum on the nonempty closed constraint set $\{f : \Phi f = y\}$. Let h be such a minimizer. Optimality of f_λ in (9.17) gives

$$\frac{1}{2\lambda} \|\Phi f_\lambda - y\|^2 + J(f_\lambda) \leq \frac{1}{2\lambda} \|\Phi h - y\|^2 + J(h) = J(h).$$

The inequality $J(f_\lambda) \leq J(h)$ and coercivity of J make $(f_\lambda)_\lambda$ bounded. Choose a convergent subsequence $f_{\lambda_k} \rightarrow f^\star$ as $k \rightarrow +\infty$. Write $m = \inf J > -\infty$. The estimate $\|\Phi f_{\lambda_k} - y\|^2 \leq 2\lambda_k(J(h) - m)$ and the limit $\lambda_k \rightarrow 0$ imply $\Phi f^\star = y$, so $f^\star$ is feasible for (9.18). Lower semicontinuity of J , applied to $J(f_{\lambda_k}) \leq J(h)$, gives $J(f^\star) \leq J(h)$. Thus $f^\star$ solves (9.18). $\square$

Full coercivity of J can be replaced by boundedness of joint sublevel sets of $J(f)$ and $\|\Phi f\|$. For seminorm penalties such as TV, this holds when $\ker(\Phi)$ intersects the nullspace of the seminorm only at zero; in particular, the measurements must detect constant images.

Quadratic Regularization. The simplest prior functionals are quadratic, and can be written as

$$J(f) = \frac{1}{2} \|Gf\|_{\mathbb{R}^K}^2 = \frac{1}{2} \langle Lf, f \rangle_{\mathbb{R}^N} \quad (9.19)$$

where $G \in \mathbb{R}^{K \times N}$ and $L = G^*G \in \mathbb{R}^{N \times N}$ is a positive semidefinite matrix. The special case (9.10) is recovered when setting $G = L = \operatorname{Id}_N$.

Writing down the first-order optimality conditions for (9.17) leads to

$$\nabla \mathcal{E}(f) = \Phi^*(\Phi f - y) + \lambda Lf = 0,$$

hence, if

$$\ker(\Phi) \cap \ker(G) = \{0\},$$

then (9.17) has a unique minimizer f_λ , which is obtained by solving a linear system

$$f_\lambda = (\Phi^* \Phi + \lambda L)^{-1} \Phi^* y. \quad (9.20)$$

If L is diagonal in a full right-singular basis of Φ , with eigenvalues α_m^2 , then, for the positive singular values,

$$\langle f_\lambda, v_m \rangle = \frac{\sigma_m}{\sigma_m^2 + \lambda \alpha_m^2} \langle y, u_m \rangle. \quad (9.21)$$

Example of convolution. Consider the convolution operator

$$\Phi f = h \star f, \quad (9.22)$$

and $G = \nabla$, a discretization of the gradient operator, for example using first-order finite differences (2.16). This corresponds to the discrete Sobolev prior introduced in Section 8.1.2. Such an operator computes, for a d -dimensional signal $f \in \mathbb{R}^N$ (for instance a 1-D signal for $d = 1$ or an image when $d = 2$), an approximation $\nabla f_n \in \mathbb{R}^d$ of the gradient vector at each sample location n . Thus typically, $\nabla : f \mapsto (\nabla f_n)_n \in \mathbb{R}^{N \times d}$ maps to d -dimensional vector fields. Then $-\nabla^* : \mathbb{R}^{N \times d} \rightarrow \mathbb{R}^N$ is a discretized divergence operator. In this case, $\Delta = -G^* G$ is a discretization of the Laplacian, which is itself a convolution operator. One then has

$$\hat{f}_{\lambda,m} = \frac{\hat{h}_m^* \hat{y}_m}{|\hat{h}_m|^2 - \lambda \hat{d}_{2,m}}, \quad (9.23)$$

where $\hat{d}_2$ is the Fourier transform of the filter d_2 corresponding to the Laplacian. For instance, in dimension 1, using first-order finite differences, the expression for $\hat{d}_{2,m}$ is given in (2.18).

9.4.1 Solving Linear Systems

If Φ and L are not simultaneously diagonalizable in the required bases, the coefficient formula (9.21) is unavailable. We instead solve (9.20), written as

$$Af = b \quad \text{where} \quad A \stackrel{\text{def.}}{=} \Phi^* \Phi + \lambda L \quad \text{and} \quad b = \Phi^* y.$$

For moderate N , up to a few thousand, direct methods solve this system at a cost of $O(N^3)$ for a generic A . Since A is symmetric positive definite, Cholesky factorization $A = BB^*$ is a natural choice, with B lower triangular. Reordering the rows and columns can help preserve sparsity in the factors of A .

For large N , iterative solvers access A only through matrix-vector products. In imaging, operator structure often makes these products much cheaper than a dense $O(N^2)$ computation. For a sufficiently sparse matrix A , the cost is typically $O(N)$. When A is symmetric positive definite, as assumed here, the conjugate gradient method solves the equivalent quadratic minimization problem

$$\min_{f \in \mathbb{R}^N} \mathcal{E}(f) \stackrel{\text{def.}}{=} Q(f) \stackrel{\text{def.}}{=} \frac{1}{2} \langle Af, f \rangle - \langle f, b \rangle \quad (9.24)$$

This objective agrees with (9.17) up to a constant. Starting from any $f^{(0)}$, conjugate gradients computes $f^{(\ell+1)}$ by minimizing over the affine space generated by all gradients obtained so far, rather than taking a single gradient step as in Section 19.1:

$$f^{(\ell+1)} \stackrel{\text{def.}}{=} \underset{f}{\operatorname{argmin}} \left\{ \mathcal{E}(f) ; f \in f^{(\ell)} + \operatorname{Span}(\nabla \mathcal{E}(f^{(0)}), \dots, \nabla \mathcal{E}(f^{(\ell)})) \right\}.$$

This subspace minimization requires only one matrix-vector product per iteration. For $\ell \geq 1$, with initialization $d^{(0)} = \nabla \mathcal{E}(f^{(0)}) = Af^{(0)} - b$, the recurrence is

$$f^{(\ell+1)} = f^{(\ell)} - \tau_\ell d^{(\ell)} \quad \text{where} \quad d^{(\ell)} = g^{(\ell)} + \frac{\|g^{(\ell)}\|^2}{\|g^{(\ell-1)}\|^2} d^{(\ell-1)} \quad \text{and} \quad \tau_\ell = \frac{\langle g^{(\ell)}, d^{(\ell)} \rangle}{\langle Ad^{(\ell)}, d^{(\ell)} \rangle} \quad (9.25)$$

where $g^{(\ell)} \stackrel{\text{def.}}{=} \nabla \mathcal{E}(f^{(\ell)}) = Af^{(\ell)} - b$. The iteration stops as soon as $g^{(\ell)} = 0$. It can also be shown that the directions $d^{(\ell)}$ are A -orthogonal (conjugate), so that after at most N iterations in exact arithmetic, the conjugate gradient method computes the unique solution $f^{(\ell)}$ of the linear system $Af = b$. In practice, it is usually used as an approximate solver, with far fewer than N iterations. Nonlinear conjugate-gradient methods extend this idea to smooth objectives, but require a line search and additional safeguards. The linear-system step-size formula and finite-termination property above do not carry over unchanged.

9.5 Non-Quadratic Regularization

9.5.1 Total Variation Regularization

Quadratic regularization (9.19) often produces a blurred reconstruction f_λ . This can poorly approximate a signal f_0 with sharp transitions, such as image edges. In the convolutional case (9.23), the restoration operator is a filter, which tends to smooth sharp features.

For the squared-norm penalty, the restored signal belongs to $\ker(\Phi)^\perp$, so missing components cannot be recovered. This statement does not hold for every quadratic penalty: a different matrix L can couple observed and unobserved components. Nonquadratic priors, particularly nonsmooth ones such as TV and ℓ^1 , encode structures such as edges and sparse coefficients that quadratic penalties do not model well.

Total variation. Total variation is a widely used nonquadratic, nonsmooth prior. For smooth functions $f : \mathbb{R}^d \mapsto \mathbb{R}$, it replaces the Sobolev or Dirichlet energy

$$J_{\text{Sob}}(f) \stackrel{\text{def.}}{=} \frac{1}{2} \int_{\mathbb{R}^d} \|\nabla f\|_{\mathbb{R}^d}^2 dx,$$

where $\nabla f(x) = (\partial_{x_1} f(x), \dots, \partial_{x_d} f(x))^\top$ is the gradient, by the (vectorial) L^1 norm of the gradient

$$J_{\text{TV}}(f) \stackrel{\text{def.}}{=} \int_{\mathbb{R}^d} \|\nabla f\|_{\mathbb{R}^d} dx.$$

Section 8.1.1 gives further background on these priors. Removing the square² from the integrand changes which functions have finite energy.

For a nontrivial bounded set Ω with a smooth boundary, $J_{\text{Sob}}(1_\Omega) = +\infty$, whereas $J_{\text{TV}}(1_\Omega) = \text{Per}(\Omega) < +\infty$. More generally, an integrable function f is of bounded variation when its distributional gradient Df is a finite vector-valued Radon measure. Its total variation is

$$J_{\text{TV}}(f) = |Df|(\mathbb{R}^d) = \sup_{\substack{h \in C_c^1(\mathbb{R}^d; \mathbb{R}^d) \\ \|h\|_\infty \leq 1}} \int_{\mathbb{R}^d} f(x) \operatorname{div} h(x) dx.$$

For a finite vector measure m , its total mass is equivalently

$$|m|(\mathbb{R}^d) = \sup_{\substack{h \in C_c(\mathbb{R}^d; \mathbb{R}^d) \\ \|h\|_\infty \leq 1}} \int_{\mathbb{R}^d} \langle h, dm \rangle.$$

The coarea formula for $f \in \text{BV}(\mathbb{R}^d)$ reads

$$J_{\text{TV}}(f) = \int_{\mathbb{R}} \text{Per}(\{f > t\}) dt.$$

For smooth functions, this agrees, for almost every level, with integrating the $(d-1)$ -dimensional Hausdorff measure of the level sets. For general BV functions, superlevel-set perimeters, rather than measures of the pointwise sets $\{f = t\}$, are required.

Discretized Total variation. For discrete data $f \in \mathbb{R}^N$, the construction in Section 8.1.2 generalizes (8.6) to any spatial dimension:

$$J_{\text{TV}}(f) = \sum_n \|\nabla f_n\|_{\mathbb{R}^d}$$

where $\nabla f_n \in \mathbb{R}^d$ approximates the spatial gradient at grid point n .

The discrete total variation prior $J_{\text{TV}}(f)$ defined in (8.6) is a convex but nondifferentiable function of f , since a term of the form $\|\nabla f_n\|$ is nondifferentiable if $\nabla f_n = 0$. Chapters 18 and 19 develop nonsmooth convex optimization methods for such functionals.

To use classical gradient descent, we replace the nonsmooth ℓ^2 norm $\|\cdot\|$ by a differentiable approximation, for example

$$\forall u \in \mathbb{R}^d, \quad \|u\|_\varepsilon \stackrel{\text{def.}}{=} \sqrt{\varepsilon^2 + \|u\|^2}.$$

This gives the smoothed TV functional already introduced in (8.12):

$$J_{\text{TV}}^\varepsilon(f) \stackrel{\text{def.}}{=} \sum_n \|\nabla f_n\|_\varepsilon$$

For fixed u , the limits as $\varepsilon \downarrow 0$ and $\varepsilon \rightarrow +\infty$ are described by

$$\|u\|_\varepsilon \xrightarrow{\varepsilon \rightarrow 0} \|u\| \quad \text{and} \quad \|u\|_\varepsilon = \varepsilon + \frac{1}{2\varepsilon} \|u\|^2 + O(\varepsilon^{-3})$$

Thus $J_{\text{TV}}^\varepsilon$ connects J_{TV} to a rescaled version of J_{Sob} , up to an additive constant.

The resulting regularized inverse problem (9.17) thus reads

$$f_\lambda \stackrel{\text{def.}}{=} \underset{f \in \mathbb{R}^N}{\operatorname{argmin}} \mathcal{E}(f) = \frac{1}{2} \|y - \Phi f\|^2 + \lambda J_{\text{TV}}^\varepsilon(f) \quad (9.26)$$

If $\ker(\Phi) \cap \ker(\nabla) = \{0\}$, the objective is coercive and strictly convex, so it has a unique minimizer. Strict convexity of $\|\cdot\|_\varepsilon$ alone does not ensure uniqueness after composition with the gradient, whose nullspace contains constants.

9.5.2 Gradient Descent Method

The optimization program (9.26) is an example of smooth unconstrained convex optimization of the form

$$\min_{f \in \mathbb{R}^N} \mathcal{E}(f) \quad (9.27)$$

where $\mathcal{E} : \mathbb{R}^N \rightarrow \mathbb{R}$ is a C^1 function. Recall that the gradient $\nabla \mathcal{E} : \mathbb{R}^N \mapsto \mathbb{R}^N$ of this functional (not to be confused with the discretized gradient $\nabla f \in \mathbb{R}^{N \times d}$ of f) is defined by the following first-order relation

$$\mathcal{E}(f + r) = \mathcal{E}(f) + \langle \nabla \mathcal{E}(f), r \rangle_{\mathbb{R}^N} + o(\|r\|_{\mathbb{R}^N})$$

A locally Lipschitz gradient improves the remainder to $O(\|r\|^2)$. Mere continuity of the gradient does not give this stronger estimate.

For such a function, the gradient descent algorithm is defined as

$$f^{(\ell+1)} \stackrel{\text{def.}}{=} f^{(\ell)} - \tau_\ell \nabla \mathcal{E}(f^{(\ell)}), \quad (9.28)$$

The step size $\tau_\ell > 0$ must balance stability against progress per iteration.

Section 19.1 analyzes convergence and explains how it depends on the step size τ_ℓ .

9.5.3 Examples of Gradient Computation

Note that the gradient of a quadratic function $Q(f)$ of the form (9.24) reads

$$\nabla Q(f) = Af - b.$$

In particular, the first-order optimality condition $\nabla Q(f) = 0$ is equivalent to the linear system $Af = b$.

For the quadratic fidelity term $\mathcal{G}(f) = \frac{1}{2} \|\Phi f - y\|^2$, one thus obtains

$$\nabla \mathcal{G}(f) = \Phi^*(\Phi f - y).$$

In the special case of the regularized TV problem (9.26), the gradient of $\mathcal{E}$ reads

$$\nabla \mathcal{E}(f) = \Phi^*(\Phi f - y) + \lambda \nabla J_{\text{TV}}^\varepsilon(f).$$

For differentiable maps $\mathcal{F} : \mathbb{R}^P \rightarrow \mathbb{R}$ and $\mathcal{H} : \mathbb{R}^N \rightarrow \mathbb{R}^P$, the chain rule gives

$$\nabla(\mathcal{F} \circ \mathcal{H})(f) = D\mathcal{H}(f)^* \nabla \mathcal{F}(\mathcal{H}(f)).$$

Apply it to $J_{\text{TV}}^\varepsilon = \|\cdot\|_{1,\varepsilon} \circ \nabla$, where $\|u\|_{1,\varepsilon} = \sum_n \sqrt{\varepsilon^2 + \|u_n\|^2}$. The result is

$$\nabla J_{\text{TV}}^\varepsilon(f) = \nabla^* \mathcal{N}^\varepsilon(\nabla f) = -\operatorname{div}(\mathcal{N}^\varepsilon(\nabla f)), \quad \mathcal{N}^\varepsilon(u)_n = \frac{u_n}{\sqrt{\varepsilon^2 + \|u_n\|^2}}.$$

Thus the gradient of the energy is computed by taking finite differences, normalizing the resulting vector field, and applying minus the divergence.

Since $\text{div} = -\nabla^*$, their operator norms are equal. For unscaled forward differences in d dimensions with periodic boundaries, $\|\nabla\|_{\text{op}} \leq 2\sqrt{d}$. The Hessian of $u \mapsto \sqrt{\varepsilon^2 + \|u\|^2}$ is

$$\frac{\text{Id}_d}{\sqrt{\varepsilon^2 + \|u\|^2}} - \frac{uu^*}{(\varepsilon^2 + \|u\|^2)^{3/2}},$$

and is bounded above by $\varepsilon^{-1}\text{Id}_d$. Consequently, a Lipschitz constant for the gradient of $\mathcal{E}$ is

$$L = \|\Phi\|_{\text{op}}^2 + \frac{\lambda \|\nabla\|_{\text{op}}^2}{\varepsilon}.$$

The smallest eigenvalue of this pointwise Hessian is $\varepsilon^2/(\varepsilon^2 + \|u\|^2)^{3/2}$, which tends to zero as $\|u\| \rightarrow \infty$. Thus smoothed TV is not globally strongly convex. If Φ is injective, the fidelity term supplies a global strong-convexity constant $\sigma_{\min}(\Phi)^2$; otherwise additional analysis on bounded sublevel sets is needed.

A fixed step $0 < \tau < 2/L$ gives convergence of gradient descent when a minimizer exists. As $\varepsilon \downarrow 0$, this bound forces small steps, motivating the nonsmooth methods developed in later chapters. A modest positive ε can also reduce the staircase artifacts of exact TV regularization.

9.6 Examples of Inverse Problems

We now discuss several inverse problems in imaging that can be solved using quadratic regularization or nonlinear TV.

9.6.1 Deconvolution

The blurring operator (9.1) is diagonal in the Fourier domain, so quadratic regularization can be solved efficiently using fast Fourier transforms under periodic boundary conditions. We refer to (9.22) and the corresponding explanations. TV regularization does not reduce to a single diagonal Fourier solve and requires an iterative optimization algorithm.

9.6.2 Inpainting

For the inpainting problem, the operator defined in (9.2) is diagonal in space

$$\Phi = \text{diag}_m(\delta_{\Omega^c}[m]),$$

and is an orthogonal projector $\Phi^* = \Phi$.

For noiseless data, we enforce the affine constraint $\{f \in \mathbb{R}^N ; y = \Phi f\}$ using the orthogonal projector

$$\forall x, \quad P_y(f)(x) = \begin{cases} f(x) & \text{if } x \in \Omega, \\ y(x) & \text{if } x \notin \Omega. \end{cases}$$

Projected gradient descent solves (9.18) for a smooth prior. For the Sobolev energy, the algorithm iterates

$$f^{(\ell+1)} = P_y(f^{(\ell)} + \tau \Delta f^{(\ell)}).$$

The iteration converges if $0 < \tau < 1/4$, since $\|\Delta\| \leq 8$. Figure 9.4 shows how it progressively fills the missing region.

Figure 9.5 illustrates Sobolev inpainting on a flower image with a known grid-shaped occlusion.

For the smoothed TV prior, the gradient descent reads

$$f^{(\ell+1)} = P_y \left(f^{(\ell)} + \tau \text{div} \left(\frac{\nabla f^{(\ell)}}{\sqrt{\varepsilon^2 + \|\nabla f^{(\ell)}\|^2}} \right) \right)$$

which converges if $0 < \tau < \varepsilon/4$.

Figure 9.6 compares Sobolev and TV inpainting on the same observed flower pixels. The displayed reconstructions illustrate their different treatment of smooth regions and edges; their SNR values are measured against the same clean image.

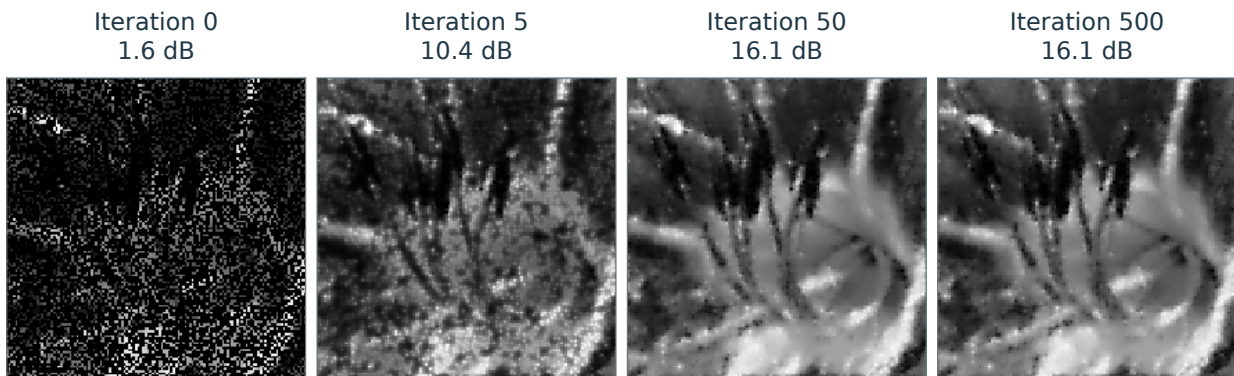


Figure 9.4. Sobolev projected gradient descent algorithm.

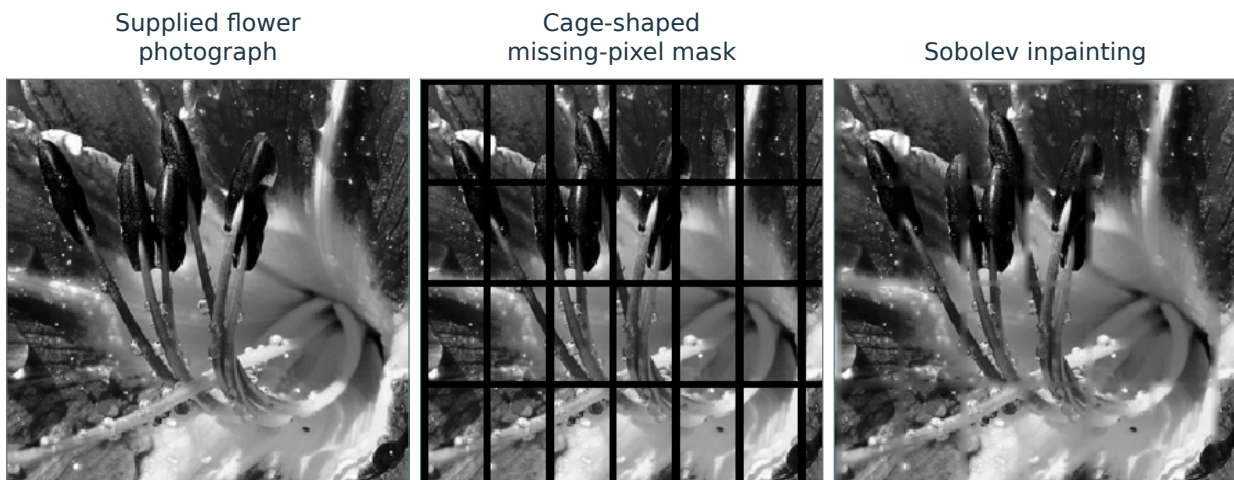


Figure 9.5. Sobolev inpainting of a grid-shaped occlusion on the flower image.

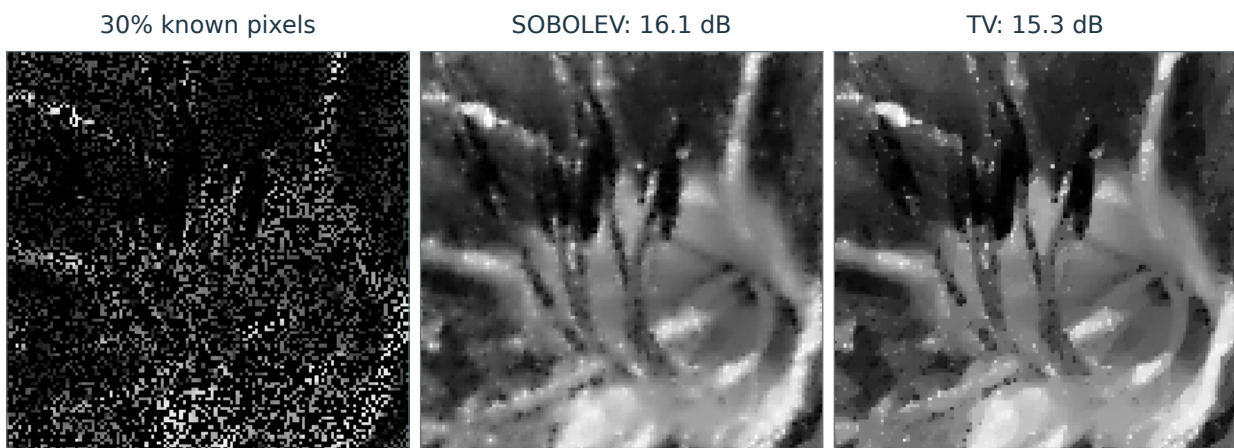


Figure 9.6. Inpainting with Sobolev and TV regularization.

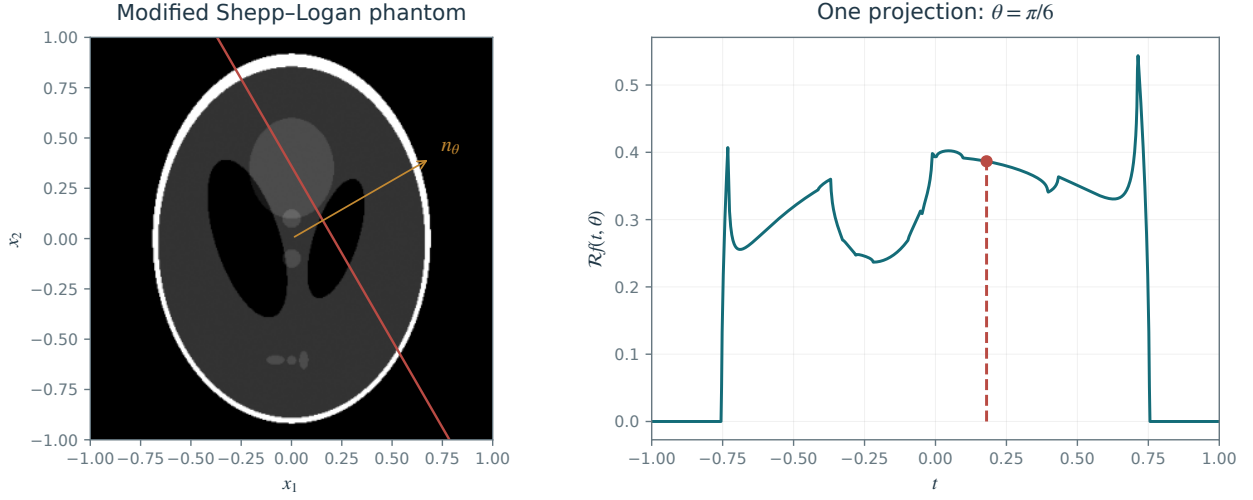


Figure 9.7. The modified Shepp-Logan phantom and one Radon projection at angle $\theta = \pi/6$. The marked line and point correspond to the same detector coordinate.

9.6.3 Tomography Inversion

In an idealized two-dimensional tomography model, each projection integrates the unknown image along a line $\Delta_{t,\theta}$ defined by

$$x \cdot \tau_\theta = x_1 \cos \theta + x_2 \sin \theta = t$$

Here $\tau_\theta = (\cos \theta, \sin \theta)$ is the unit normal to the line.

The resulting Radon transform is

$$\forall \theta \in [0, \pi), \forall t \in \mathbb{R}, \quad p_\theta(t) = \int_{\Delta_{t,\theta}} f(x) ds = \iint f(x) \delta(x \cdot \tau_\theta - t) dx$$

see Figure 9.7.

With $\hat{f}(\omega) = \int f(x) e^{-ix \cdot \omega} dx$, the Fourier slice theorem identifies the one-dimensional Fourier transform of each projection with a radial slice of the two-dimensional Fourier transform

$$\forall \theta \in [0, \pi), \forall \xi \in \mathbb{R} \quad \hat{p}_\theta(\xi) = \hat{f}(\xi \cos \theta, \xi \sin \theta). \quad (9.29)$$

Changing to polar coordinates in Fourier inversion gives the filtered-backprojection formula

$$f(x) = \frac{1}{2\pi} \int_0^\pi (p_\theta \star h)(x \cdot \tau_\theta) d\theta$$

with $\hat{h}(\xi) = |\xi|$.

We model acquisition at finitely many equally spaced orientations $\{\theta_k = \pi k/K\}_{0 \leq k < K}$ by the partial Radon transform

$$Rf = (p_{\theta_k})_{0 \leq k < K}.$$

By (9.29), knowing Rf is equivalent to knowing the Fourier transform of f along the corresponding radial lines:

$$\{\hat{f}(\xi \cos(\theta_k), \xi \sin(\theta_k))\}_k.$$

For a simplified discrete model, we therefore treat acquisition as direct sampling of the Fourier transform:

$$\Phi f = (\hat{f}[\omega])_{\omega \in \Omega} \in \mathbb{C}^P$$

where Ω consists of discrete radial sampling lines in the Fourier plane; see Figure 9.8, right.

In this idealized discrete model, recovery is an inpainting problem in the Fourier domain. We use the unitary discrete Fourier transform and consider measurements

$$\forall \omega \in \Omega, \quad y[\omega] = \hat{f}[\omega] + w[\omega]$$

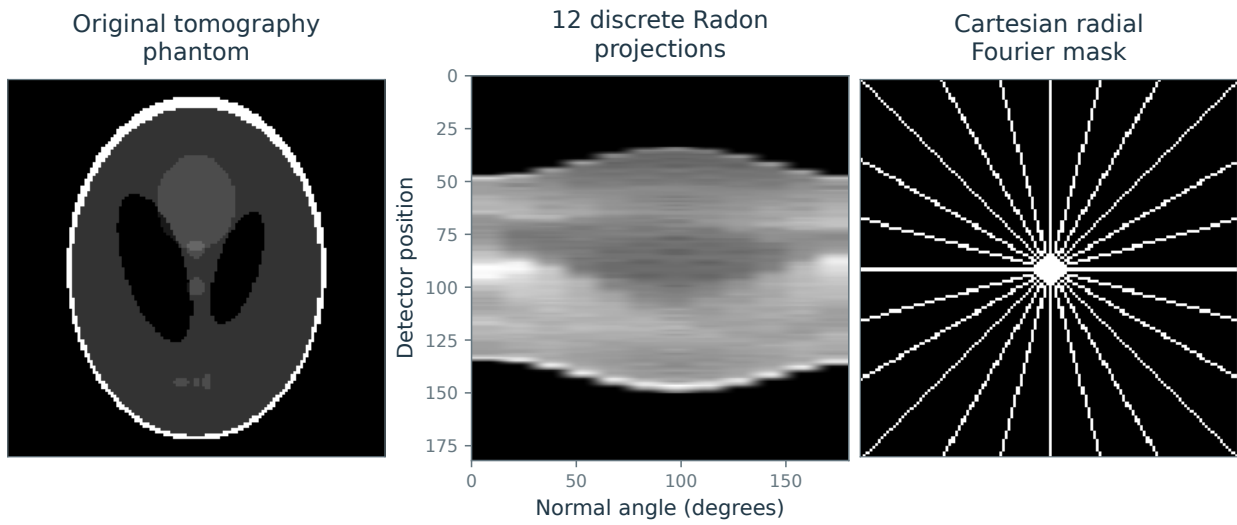


Figure 9.8. Tomography test phantom, discrete Radon projections, and a Cartesian radial Fourier-sampling mask.

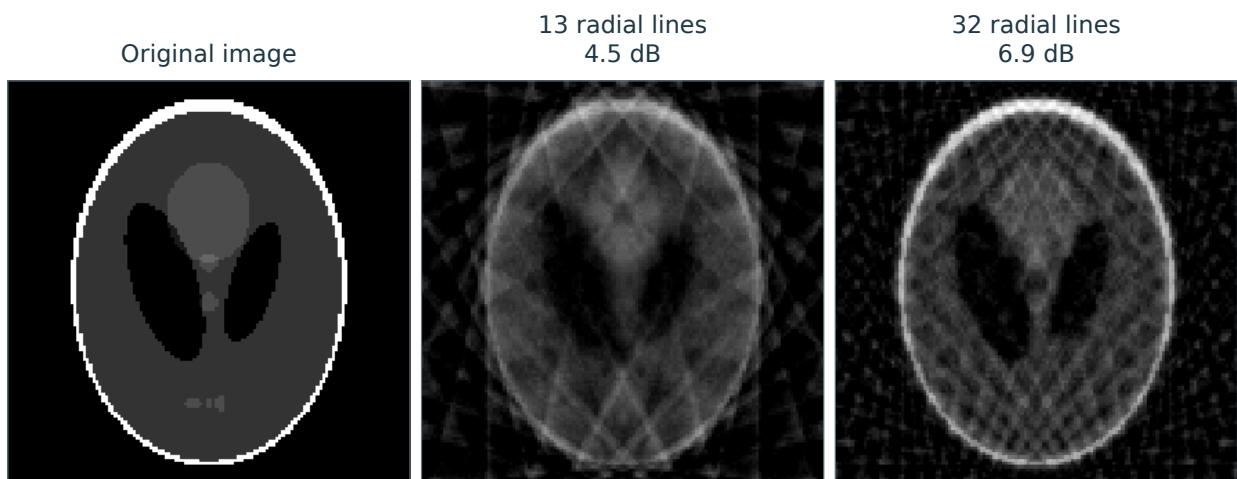


Figure 9.9. Pseudoinverse reconstruction from partial Cartesian Fourier measurements on 13 and 32 radial lines.



Figure 9.10. TV-regularized reconstruction and the pseudoinverse for the same partial Fourier measurements.

where $w[\omega]$ denotes measurement noise, modeled here as Gaussian white noise.

For these partial Fourier measurements, the pseudoinverse $f^+ = \Phi^+ y$ from (9.7) is given by

$$\hat{f}^+[\omega] = \begin{cases} y[\omega] & \text{if } \omega \in \Omega, \\ 0 & \text{if } \omega \notin \Omega. \end{cases}$$

Figure 9.9 shows examples of pseudoinverse reconstruction as the size of Ω increases. This reconstruction exhibits pronounced artifacts because missing Fourier frequencies are set to zero.

The total variation regularization (9.17) reads

$$f^\star \in \operatorname{argmin}_f \frac{1}{2} \sum_{\omega \in \Omega} |y[\omega] - \hat{f}[\omega]|^2 + \lambda \|f\|_{\text{TV}}.$$

TV regularization suits images with approximately constant regions separated by sharp boundaries, as in the cartoon model of Section 5.2.4. Figure 9.10 compares TV reconstruction and the pseudoinverse on a synthetic image of this type. The example illustrates recovery of sharp features from incomplete Fourier data. Spatial inpainting can be more challenging for TV, as Figure 9.6 illustrates.

References

- [1] Heinz Werner Engl, Martin Hanke, and Andreas Neubauer. *Regularization of inverse problems*, volume 375. Springer Science & Business Media, 1996.
- [2] Stephane Mallat. *A wavelet tour of signal processing: the sparse way*. Academic press, 2008.
- [3] Otmar Scherzer, Markus Grasmair, Harald Grossauer, Markus Haltmeier, Frank Lenzen, and L Sirovich. *Variational methods in imaging*. Springer, 2009.