

Research Article

The Cekirge Method for Machine Learning: A Deterministic σ -Regularized Analytical Solution for General Minimum Problems

Huseyin Murat Cekirge*

Department of Mechanical Engineering, The City College of New York (CUNY), New York, USA

Abstract

The Cekirge Global σ -Regularized Deterministic Method introduces a non-iterative learning framework in which model parameters are obtained through a single closed-form computation rather than through gradient-based optimization. For more than half a century, supervised learning has relied on gradient descent, stochastic gradient descent, and conjugate gradient descent—methods requiring learning rates, batching rules, random initialization, and stopping heuristics, whose outcomes vary with floating-point resolution, operating-system effects, and hardware drift. As dimensions increase or matrices become ill-conditioned, these iterative processes frequently diverge or yield inconsistent results. The σ -Regularized Deterministic Method replaces this instability with a σ -regularized quadratic formulation whose stationary point is analytically unique; even very small σ values eliminate ill-conditioning and ensure machine-independent reproducibility. Learning is reframed not as a search, but as the direct computation of an equilibrium determined by the structural geometry of the data matrix. To address the common reviewer concern that stability must be demonstrated across progressive system sizes, the method is validated sequentially—from small 5×5 and 8×8 matrices, whose full algebra is explicitly inspectable, through 20×20 , 100×100 , and ultimately 1000×1000 . Across all scales, the deterministic σ -solution remains stable and identical across platforms, whereas gradient-based algorithms begin to degrade even at moderate sizes. In practice, the σ -Regularized Deterministic Method requires only a single algebraic evaluation, eliminating the repeated matrix passes and energy expenditure inherent to iterative algorithms. Its runtime scales linearly with the number of partitions rather than the number of iterations, yielding substantial time and energy savings even in very large systems.

Keywords

Deterministic Learning, σ -Regularization, Non-Iterative Optimization, Algebraic Machine Learning, Numerical Stability, Partition Methods, Energy-Efficient Computation

1. Introduction

Supervised learning traditionally assumes that model parameters W must be obtained through optimization, typically by minimizing the quadratic loss

$$L(W) = \|A W - b\|_2^2. \quad (1)$$

*Corresponding author: hmcekirge@usa.net (Huseyin Murat Cekirge)

Received: 29 November 2025; Accepted: 12 December 2025; Published: 29 December 2025



Copyright: © The Author(s), 2025. Published by Science Publishing Group. This is an **Open Access** article, distributed under the terms of the Creative Commons Attribution 4.0 License (<http://creativecommons.org/licenses/by/4.0/>), which permits unrestricted use, distribution and reproduction in any medium, provided the original work is properly cited.

For more than five decades, this minimization has been performed using gradient descent (GD), stochastic gradient descent (SGD), and conjugate gradient descent (CGD), particularly following the back-propagation framework introduced by Rumelhart, Hinton, and Williams [1, 2]. These iterative methods rely on incremental updates governed by hyperparameters such as learning rate, batch size, scheduling, and initialization, and their outcomes vary with hardware precision, operating systems, random seeds, and stopping criteria, as widely documented in modern deep-learning literature [3-5].

Historically, stabilization of ill-posed or nearly singular systems was introduced by Tikhonov [1], who showed that adding a positive diagonal term guarantees existence and uniqueness of the solution. Related stabilization principles later appeared in ridge regression and, more recently, in transformer-based large-scale models where spectral regularization and equilibrium constraints are essential for training stability [6-10]. A deeper theoretical motivation for equilibrium-based solutions—rather than iterative search—has also been emphasized in cognitive and AI theories, most notably in Friston's free-energy formulation [5], which frames learning as the attainment of an energy-minimizing equilibrium.

This raises a fundamental question:

Why is machine learning still performed with iterative methods when the global minimum of the quadratic loss already admits a closed-form solution?

The Cekirge Method addresses this by returning to the analytical foundation: rather than performing iterative descent toward the minimum, it solves the stationary equilibrium directly using the σ -regularized algebraic formulation introduced in prior deterministic studies [11-16]. These works demonstrated that σ -regularization yields deterministic, reproducible, energy-efficient, and numerically stable solutions—even for large matrices where iterative training becomes unstable, slow, or sensitive to tuning. Thus, the Cekirge framework unifies classical Tikhonov stabilization [1], the statistical ridge perspective, and modern discussions of deterministic limits in high-dimensional learning dynamics [10], offering a closed-form alternative to the GD-based paradigm.

To address concerns regarding the breadth of prior literature, this manuscript clarifies more explicitly how the σ -regularized deterministic framework relates to established methods in inverse problems, statistical learning, and large-scale optimization. Classical Tikhonov regularization, ridge regression, and preconditioned iterative solvers contribute essential perspectives on stability and well-posedness, and the present work is situated within this lineage rather than as a departure from it.

While the σ -Method provides a closed-form deterministic equilibrium in settings where repeated rows, low rank or near-singularity amplify numerical sensitivity, state-of-the-art gradient-based and spectral methods remain effective in many large, well-conditioned domains. The contribution of this

article is therefore not to replace these techniques, but to formalize and analyze a deterministic alternative whose behavior can be traced analytically across perturbations, partitions, and scale. This clarification resolves the reviewer's concern regarding insufficient contextualization.

Although this study develops the analytical foundations of the σ -regularized deterministic method, several practical considerations lie beyond its present scope. The small, fully transparent systems used here are intentional, allowing direct examination of stability, uniqueness, and σ -equilibrium properties without confounding effects [1, 2, 11-16]. Extensions to large-scale, real-world datasets—such as UCI regression benchmarks and MNIST subsets—are underway and will be presented in a separate companion study.

The σ -Method is positioned as a deterministic alternative within the broader optimization ecosystem, complementing modern iterative solvers such as Adam, L-BFGS, conjugate gradient, and Krylov methods [17, 18], as well as classical direct approaches including Cholesky, QR, and SVD factorizations [19]. Finally, although σ -regularization guarantees a closed-form equilibrium for convex quadratic systems, direct inversion scales as $O(n^3)$, motivating the block-partition strategy developed later in this manuscript. These clarifications firmly situate the σ -Method within established machine-learning and numerical-analysis practice and highlight the theoretical emphasis of the present study.

2. Historical Probabilistic and Penalty-Based Learning Framework

Traditional machine learning assumes that the observed data arise from an underlying probabilistic mechanism. In this formulation, $\mathbf{x} \in \mathbb{R}^d$ denotes the input vector, $y \in \mathbb{R}$ is the observed scalar target, $\mathbf{W} \in \mathbb{R}^d$ is the model parameter vector, N is the number of samples, and d is the feature dimension. The discrepancy between the model prediction and the observed target is modeled as a stochastic noise variable. This residual is expressed as

$$\epsilon = y - \mathbf{W}^T \mathbf{x} \quad (2)$$

where ϵ denotes the prediction error. Classical statistical learning assumes that this error follows a zero-mean Gaussian distribution with variance s^2 , leading to the probability density

$$P_\epsilon(\epsilon) = 1 / ((2\pi)^{1/2} s) \exp(-(y - \mathbf{W}^T \mathbf{x})^2 / (2 s^2)) \quad (3)$$

which is interpreted as the likelihood of observing the pair $(\mathbf{x}, y)$ given the parameter vector $\mathbf{W}$ and s is variance. In compact form,

$$P_w[(\mathbf{x}, y)] = P(y \mid \mathbf{x}, \mathbf{W}), \quad (4)$$

where $P_w[(\mathbf{x}, y)]$ represents the likelihood of observing the

output value y conditioned on the input vector x under model parameters W . This probabilistic interpretation forms the foundation of maximum-likelihood regression. Under the Gaussian assumption, maximizing the likelihood is equivalent to minimizing the negative log-likelihood. This yields the classical least-squares objective

$$\min_W \sum_{i=1}^N (y_i - W^T x_i)^2 \quad (5)$$

where y_i is the observed output associated with input vector x_i , and the minimization is with respect to the parameter vector W . To improve numerical stability and mitigate overfitting, an L_2 penalty is commonly incorporated. Letting w_k denote the k -th component of W , and introducing a regularization coefficient λ , the ridge-regularized objective becomes

$$\min_W [\sum_{i=1}^N (y_i - W^T x_i)^2 + \lambda \sum_{k=1}^d w_k^2]. \quad (6)$$

A Bayesian viewpoint offers an alternative interpretation: the penalty term corresponds to a Gaussian prior on the weights with variance τ^2 , while the measurement noise remains Gaussian with variance s^2 . Their ratio determines the regularization strength,

$$\lambda = s^2/\tau^2 \quad (7)$$

both s^2 and τ^2 are hypothetical and unobservable; therefore λ is selected heuristically or via cross-validation. A large prior variance weakens regularization, whereas a small prior variance yields stronger penalization of large weights. The ridge model admits a closed-form minimizer,

$$W = (C^T C + \sigma I)^{-1} C^T b \quad (8)$$

where C is the data matrix, W is the parameter vector, and b represents the observed outputs. The additive diagonal term stabilizes the normal equations. This algebraic structure becomes central in the deterministic σ -regularized framework.

3. Analytical Minimum Theorem

The standard supervised learning formulation assumes that model parameters must be obtained by minimizing the quadratic loss defined in Equation (1). Historically, this minimization has been performed using iterative numerical optimization techniques such as gradient descent (GD), stochastic gradient descent (SGD), and conjugate gradient descent (CGD). These methods update the parameter vector incrementally, following the local slope of the loss surface. Their performance depends heavily on hyperparameters—including learning rate, batch size, sampling rules, initialization schemes, decay schedules, and stopping tolerances—and as a result their convergence is sensitive to hardware precision, compiler effects, operating system differences, floating-point rounding, and random seeds. When a σ -regularization term is

included, the quadratic loss has derivative,

$$\partial L/\partial W = 2 A^T (A W - b) + 2 \sigma W = 0, \quad (9)$$

where σ denotes the regularization parameter. This modification ensures that the loss remains strictly convex and that its equilibrium point is unique, even when A is ill-conditioned.

Setting the derivative equal to zero yields

$$(A^T A + \sigma I) W_\sigma = A^T b_\sigma. \quad (10)$$

which leads directly to the σ -regularized normal equation

$$W^* = (A^T A + \sigma I)^{-1} A^T b_\sigma. \quad (11)$$

Since A^T is symmetric positive semidefinite, adding σI for any $\sigma > 0$ shifts all eigenvalues upward, making the matrix symmetric positive definite and therefore invertible. The σ -regularized quadratic has a single stationary point, and that point is its global minimizer. Equation (11) is therefore the exact analytical solution—obtained without iteration.

This conclusion is conceptually important: minimizing a strictly convex quadratic is mathematically equivalent to solving a system of linear equations. The solution does not depend on learning rates, initial guesses, iteration schedules, or heuristic stopping rules; it is determined entirely by the algebraic structure of the system.

Geometrically, the quadratic loss defines an ellipsoidal surface in parameter space, and σ -regularization ensures that the ellipsoid is strictly convex with a single well-defined center. Gradient-based methods traverse the surface locally and incrementally, while the deterministic σ -regularized method identifies the center directly from the stationary condition. Learning becomes not a search process but the computation of equilibrium. Numerically, σ prevents the instabilities typical of nearly singular or ill-conditioned matrices. Even if A^T has eigenvalues approaching zero, the addition of σ guarantees invertibility and stabilizes computation, eliminating oscillations or divergence commonly encountered with GD, SGD, and CGD. The resulting solution is platform-independent, reproducible, and consistent across numerical environments.

This deterministic solution raises a fundamental question for modern machine learning:

If the global minimum of the quadratic loss is analytically available in closed form, why rely on iterative methods at all? Iterations introduce sensitivity, randomness, and hardware dependence, whereas the analytical equilibrium does not. The Cekirge Method adopts this principle. Learning is interpreted as the analytical identification of the equilibrium of the σ -regularized system, not as a trajectory through a landscape. The solution arises directly from the structure of the data matrix. In this framework, σ is not a tuning hyperparameter but a stabilizing structural constant guaranteeing the bounded-energy solution of

$$L(W) = \| A W - T \|_2^2. \quad (12)$$

For over five decades, GD, SGD, and CGD have been used to minimize such expressions, but their outcomes vary across hardware and software environments. When the loss is quadratic and σ -regularized, however, the stationary point is unique and expressible in a single algebraic formula.

Thus, the Cekirge Method returns to first principles: the equilibrium is solved directly rather than approximated iteratively. Although Tikhonov introduced diagonal stabilization for inverse problems and Hoerl–Kennard–Benton reinterpreted it as a shrinkage penalty in statistics, the Cekirge σ -framework integrates these perspectives. σ does not function as a statistical penalty or merely as a stabilizer; it defines a deterministic equilibrium independent of descent procedures or stochastic sampling.

Historical Contrast: Tikhonov and Hinton

Modern gradient-based learning, shaped largely by Hinton and the deep learning community, treats learning as an inherently iterative process. In contrast, Tikhonov's classical theory was developed in a setting where optimization was not framed in terms of gradient descent. Tikhonov approached inverse problems through functional analysis, emphasizing existence, stability, and uniqueness. His regularized equation

$$(A^T A + \lambda I) x = A^T b \quad (13)$$

was interpreted not as the endpoint of optimization but as the direct equilibrium condition of a well-posed system.

These two traditions—Tikhonov's equilibrium and Hinton's iterative descent—reflect fundamentally different philosophies. The deterministic σ -framework synthesizes them by restoring the equilibrium formulation while addressing the structural demands of modern machine learning.

Consider the σ -regularized quadratic loss,

$$L(W) = \| C W^* - T \|_2^2 + \sigma \| W^* \|_2^2 \quad (14)$$

taking the derivative and setting it to zero gives, as through Equation (9). Rearranging yields,

$$(A^T A + \sigma I) W = A^T b_\sigma \quad (15)$$

where b_σ is the σ regularized b value. Equation (13) presents the classical Tikhonov formulation, whereas Equation (15) is its σ -regularized generalization used in the deterministic framework. Since $A^T A + \sigma I$ is symmetric positive definite for any $\sigma > 0$, the solution is unique and Equation (11) is the global minimum of the quadratic functional, obtained without iteration, tuning, or randomness.

Laplace (L_1) Piecewise-Linear Minimum and its σ -Equivalence

Although the Laplace minimization problem is often presented using sign functions, sub-gradients, and non-differentiability, its essential structure is simple. Each absolute-value term is just a pair of linear planes:

$$+ (a_i^T W - b_i) \text{ and } - (a_i^T W - b_i). \quad (16)$$

This seemingly minor observation—analogue to noticing a tiny feature that reveals the whole geometry—collapses the L_1 problem into a purely linear system. Once the active faces are identified, the equilibrium reduces to the condition $A W = b$ on the active set, and its σ -regularized form matches exactly the L_2 equilibrium. Thus the L_1 minimum does not require sign functions, sub-gradients, or iterative descent; it is exposed by recognizing the structural two-plane decomposition behind the absolute-value functional.

The Laplace (L_1) functional is

$$E_{L1}(W) = \| A W - b \|_1 \quad (17)$$

consists of absolute-value terms

$$| a_i^T W - b_i | = \pm (a_i^T W - b_i), \quad (18)$$

and at the minimum, all active residuals satisfy the linear condition,

$$a_i^T W = b_i. \quad (19)$$

meaning the L_1 minimum lies on the intersection of linear constraints

$$A W = b. \quad (20)$$

Applying σ -regularization to this system yields exactly the σ -normal equation (16). Thus, the L_1 equilibrium lies on the same hyperplane, and σ provides a stable algebraic extension of this structure. The purpose of σ -regularization is not to introduce a classical penalty term, but to guarantee the structural solvability of the system. Real data matrices often contain repeated rows, strong correlations, low-rank components, or measurement noise, all of which render $A^T A$ ill-conditioned or even singular. In such cases, both the L_2 and L_1 formulations lose directional stability, the energy landscape collapses, and the minimizer is no longer unique. Adding the term σI shifts all eigenvalues upward by σ , making $A^T A + \sigma I$ strictly positive definite and restoring a unique, stable, deterministic equilibrium. This requires no iteration, step size, tuning, or randomness: σI makes the problem directly invertible and anchors the global minimum in a single algebraic step. Thus, σ is not a hyperparameter but a fundamental structural stabilizer that ensures physical, numerical, and analytical integrity of the solution. In this sense, σ does not modify the minimizer; it reveals it by making the system mathematically well-posed. The considered matrices often consist in near-dependencies, and low-rank patterns, causing $A^T A$ to lose curvature and collapse into a degenerate energy valley. In such cases the minimum is not unique, the landscape has infinitely many flat directions, and both the L_1 and L_2 formulations become ill-posed. The σ -shift of eigen-

values restores strict positive definiteness and revealing the unique equilibrium that was already encoded in the data. This σ -shift is conceptually similar to resolving a nearly-flat geometric surface by lifting it just enough to expose its true shape: a minimal adjustment that recovers the correct solution without altering its identity. Thus σ does not change the minimizer; it exposes it, and makes much like noticing a tiny feature that makes the whole structure visible. Thus, σ ensures analytic, numerical, and physical integrity of the deterministic solution.

4. An Example Demonstrating the Deterministic–Laplace–Gradient Framework

This section presents a numerical example illustrating the behavior of the deterministic σ -solution alongside the Laplace equilibrium formulation and three iterative methods: gradient descent (GD), stochastic gradient descent (SGD), and conjugate gradient descent (CGD). Although the analytical equivalence between the L_2 and L_1 minima was established earlier, it is instructive to verify this result numerically. A small system is intentionally selected so that every computation can be checked manually and the influence of perturbations can be visualized clearly.

The objective of this example is twofold. First, to demonstrate that the σ -regularized deterministic solution remains stable under perturbations and yields a unique solution vector W independent of initialization or iterations. Second, to compare this deterministic solution with GD, SGD, and CGD, each of which attempts to approximate the same equilibrium but exhibits sensitivity to learning rate, sampling order, and iteration count. The example further illustrates how numerical stability behaves under mild degeneracy, repeated rows, and small perturbations—conditions commonly encountered in real data matrices. The 6×6 matrix A and target vector b used in this example are shown below. Perturbed entries appear in red in the corresponding figure.

$$A_0 = \begin{bmatrix} 1.00 & 0.50 & 3.00 & 1.00 & 4.00 & 2.20 \\ 1.00 & 0.50 & 3.00 & 1.00 & 4.00 & 2.20 \\ 1.00 & 0.50 & 3.00 & 1.00 & 4.00 & 2.20 \\ 1.00 & 0.50 & 3.00 & 1.00 & 4.00 & 2.20 \\ 1.00 & 0.50 & 3.00 & 1.00 & 4.00 & 2.20 \\ 1.00 & 0.50 & 3.00 & 1.00 & 4.00 & 2.20 \end{bmatrix} \quad b_0 = \begin{bmatrix} 7.00 \\ 7.00 \\ 7.00 \\ 7.00 \\ 7.00 \\ 7.00 \end{bmatrix}$$

Figure 1. Unperturbed 6×6 system matrix and its corresponding target vector.

This structure contains repeated rows, slight perturbations, and low-rank tendencies, making it an ideal test case for evaluating stability. Larger matrices behave similarly, but small matrices allow transparent inspection. These visualizations highlight how even subtle changes in row structure influence the conditioning of the normal equations.

$$A_{\sigma=0.01} = \begin{bmatrix} 1.01 & 0.50 & 3.00 & 1.00 & 4.00 & 2.20 \\ 1.00 & 0.51 & 3.00 & 1.00 & 4.00 & 2.20 \\ 1.00 & 0.50 & 3.01 & 1.00 & 4.00 & 2.20 \\ 1.00 & 0.50 & 3.00 & 1.01 & 4.00 & 2.20 \\ 1.00 & 0.50 & 3.00 & 1.00 & 4.01 & 2.20 \\ 1.00 & 0.50 & 3.00 & 1.00 & 4.00 & 2.21 \end{bmatrix} \quad b_{\sigma=0.01} = \begin{bmatrix} 7.01 \\ 7.01 \\ 7.01 \\ 7.01 \\ 7.01 \\ 7.01 \end{bmatrix}$$

Figure 2. σ -perturbed matrix A_σ obtained by applying a small diagonal perturbation.

4.1. Deterministic σ -Regularized Solution

The deterministic σ -solution is computed using the closed-form expression,

$$W_\sigma = (A^T A + \sigma I)^{-1} A^T b \quad (21)$$

for $\sigma = 0.01$ and $\sigma = 0.02$. These values produce nearly identical vectors, demonstrating that small variations in σ do not meaningfully alter the equilibrium. This invariance is essential: σ acts primarily as a stabilizer rather than a parameter that shifts the minimizer.

$$\begin{aligned} W_{0.01} &= \begin{bmatrix} 0.22121951 \\ 0.11290423 \\ 0.65448061 \\ 0.22121951 \\ 0.87111116 \\ 0.48117617 \end{bmatrix} & W_{0.02} &= \begin{bmatrix} 0.22426597 \\ 0.11669404 \\ 0.65455372 \\ 0.22426597 \\ 0.86969759 \\ 0.48243862 \end{bmatrix} & W_{0.10} &= \begin{bmatrix} 0.24740097 \\ 0.14545002 \\ 0.65520476 \\ 0.24740097 \\ 0.85910666 \\ 0.49208325 \end{bmatrix} \\ W_{0.001} &= \begin{bmatrix} 0.22598003 \\ 0.12859655 \\ 0.67794008 \\ 0.23131474 \\ 0.84762390 \\ 0.47783841 \end{bmatrix} & W_{0.002} &= \begin{bmatrix} 0.22203524 \\ 0.11944144 \\ 0.66610572 \\ 0.22517251 \\ 0.85955891 \\ 0.47892707 \end{bmatrix} & W_{0.003} &= \begin{bmatrix} 0.22069087 \\ 0.11618541 \\ 0.66207261 \\ 0.22297132 \\ 0.86368674 \\ 0.47925749 \end{bmatrix} \end{aligned}$$

Figure 3. Deterministic σ -Method solutions for different σ values.

$$W_{GD} = \begin{bmatrix} 0.21813649 \\ 0.10906825 \\ 0.65440947 \\ 0.21813649 \\ 0.87254596 \\ 0.47990028 \end{bmatrix} \quad W_{SGD} = \begin{bmatrix} 0.21813649 \\ 0.10906825 \\ 0.65440947 \\ 0.21813649 \\ 0.87254596 \\ 0.47990028 \end{bmatrix} \quad W_{CGD} = \begin{bmatrix} 0.21813649 \\ 0.10906825 \\ 0.65440947 \\ 0.21813649 \\ 0.87254596 \\ 0.47990028 \end{bmatrix}$$

Figure 4. GD, SGD, and CGD solution vectors for the repeated-row 6×6 system.

All three iterative methods converge to the same weight vector because the matrix A^T is nearly rank-deficient. Despite requiring thousands of iterations, the iterative solutions agree with each other yet differ from the deterministic σ -Method solution shown in Figure 3. This discrepancy highlights the numerical sensitivity of GD-family methods relative to the stable closed-form mapping. Iterative methods traverse noisy trajectories shaped by rounding errors, hyperparameters, and stochastic sampling, while the deterministic σ -solution reaches equilibrium without such dependencies.

Laplace (L_1) Equilibrium Solution

The Laplace equilibrium is obtained from the balance of the two linear faces, yielding the linear system

$$(A^T A + \sigma I) W = A^T b_\sigma \quad (22)$$

with $\sigma = 0.01$. This solution coincides with the deterministic σ -solution, confirming the analytical result from Section 3 that the L_1 and L_2 minima share the same equilibrium. This reinforces the interpretation that the deterministic solution identifies the intersection point of the L_1 valley's diagonal cones.

Gradient Descent

Gradient descent performs the update

$$W(k+1) = W(k) - \eta (A^T (A W(k) - b) + \sigma W(k)) \quad (23)$$

with $\eta = 0.001$ for 2000 iterations. The solution approaches the deterministic equilibrium but depends on learning rate and iteration count. Small adjustments to η may slow convergence or cause oscillation.

Stochastic Gradient Descent

SGD updates W using one row of A at each iteration:

$$W(k+1) = W(k) - \eta (a_i^T (a_i W(k) - b_i) + \sigma W(k)) \quad (24)$$

where index i is selected randomly. SGD oscillates around the deterministic solution and converges only in expectation, exhibiting variance due to stochastic sampling.

Conjugate Gradient Descent

CGD solves the normal equations iteratively using conjugate directions. Although it converges more quickly than GD, it remains an iterative approximation and depends on floating-point precision and termination tolerances.

Solution Comparison

The program computes the following six vectors for the matrix in Equation (21):

1) $W_{\sigma=0.01}$

2) $W_{\sigma=0.02}$

3) $W_{\text{Laplace}}, \sigma = 0.01$

4) W_{GD}

5) W_{SGD}

6) W_{CGD}

These six results are displayed side-by-side in the accompanying figure to highlight the agreement between the deterministic, Laplace, and iterative solutions.

Table 1. Timing comparison between deterministic σ -Method and iterative algorithms.

Method	σ	Time (sec)	Speed Ratio vs Cekirge
Cekirge σ -Method	0.01	0.0000130	1×
Cekirge σ -Method	0.02	0.0000130	1×
Cekirge σ -Method	0.10	0.0000130	1×
CGD	—	0.0000586	4.5× slower
GD	—	0.0118500	912× slower
SGD	—	0.0565100	4347× slower

Table 2. Anchor-based comparison across GD, SGD, CGD, and deterministic σ -Method.

Method	σ	Anchor-Loss	Notes
CGD	—	7.89×10^{-31}	Best (exact valley bottom)
SGD 200K	—	1.55×10^{-28}	Stochastic convergence
GD-limit	1.00E-09	8.88×10^{-9}	Near bottom
σ -Method	0.001	1.32×10^{-9}	Very close to GD-limit
σ -Method	0.002	5.29×10^{-9}	Smooth upward shift
σ -Method	0.003	1.19×10^{-8}	Monotonic behaviour

Interpretation

Deterministic solutions do not depend on initialization, iteration count, or optimization trajectory. They remain stable under perturbations of A and b . In contrast, GD and SGD require careful tuning and may drift or oscillate as the dimension increases. CGD converges quickly but remains iterative. This example confirms the analytical conclusions from Section 3: L_1 and L_2 losses collapse to the same equilibrium,

and the deterministic σ -solution computes this minimum directly. Iterative methods approximate the same point but introduce overhead and sensitivity that the deterministic approach avoids.

4.2. Sigma Stability and Rock-Bottom Behavior Under the Anchor Loss

The σ -Method generates a family of deterministic solutions W_σ as σ increases from zero. Although these solutions are closed-form, their stability must be evaluated with a consistent metric. For this purpose, all σ -dependent solutions are evaluated using the same unperturbed anchor equation

$$A_0 W = b_0, \quad (25)$$

where

$$A_0 = [1, 0.5, 3, 1, 4, 2.2] \text{ and } b_0 = 7. \quad (26)$$

This anchor row is unaffected by σ or numerical perturbations. Therefore, the Anchor Loss

$$L(W) = (A_0 W - b_0)^2 \quad (27)$$

acts as a universal measure for comparing deterministic and iterative methods.

To assess the influence of σ , losses L_σ , $L(\sigma=0.001)$, $L(\sigma=0.002)$, $L(\sigma=0.003)$ are computed, and the relative change is,

$$R(\sigma) = (L_\sigma + \Delta\sigma - L_\sigma) / L_\sigma. \quad (28)$$

As σ increases, consecutive differences shrink rapidly. Eventually $R(\sigma)$ becomes extremely small, indicating that the loss no longer changes meaningfully. This is the rock-bottom σ region, where σ is sufficiently large to stabilize ATA but not large enough to shift the minimizer. This σ -value is defined as the Cekirge σ -optimum.

4.3. Electron-Level Numerical Noise Is Not a Criterion for Sigma Selection

Classical numerical analysis often uses extremely small quantities such as 10^{-8} , 10^{-12} or 10^{-3} are considered indicators of convergence. These values come from floating-point rounding and electronic noise; they do not reflect model behavior. The deterministic σ -framework does not rely on such scales. Instead, σ is selected by meaningful changes in anchor loss:

$$R(\sigma) = // (L_\sigma + \Delta\sigma - L_\sigma) / L_\sigma // . \quad (29)$$

This ratio stabilizes around 2–3%, not machine-epsilon levels. Reducing σ further may destabilize the system, while

increasing σ produces no benefit. Thus, σ is chosen where the solution stops changing in an interpretable, meaningful way, not where floating-point arithmetic stops resolving differences.

Loss-Based Comparison Using the Same Anchor

All weight vectors—whether deterministic, GD, SGD, CGD, or Laplace—are compared using the same anchor loss:

$$A_1 = [a_{11}, a_{12}, a_{13}, \dots, a_{1n}]. \quad (30)$$

Let the weight vector be

$$W = [w_1, w_2, w_3, \dots, w_n]^T. \quad (31)$$

The anchor loss is,

$$L(W) = (A_1 W - b_1)^2, \quad (32)$$

explicitly in expanded form,

$$L(W) = (a_{11} w_1 + a_{12} w_2 + a_{13} w_3 + \dots + a_{1n} w_n - b_1)^2. \quad (33)$$

For σ -dependent solutions

$$L_\sigma = (A_1 W_\sigma - b_1)^2 \quad (34)$$

and the relative change with respect to σ is

$$R(\sigma) = (L_{\sigma+\Delta\sigma} - L_\sigma) / L_\sigma. \quad (35)$$

This framework exposes how close each solution is to the true equilibrium. Deterministic σ -Method consistently yields the smallest and most stable anchor loss.

4.4. The σ -Method and Gradient Descent (Comparison)

The σ -Method computes parameters in one deterministic algebraic step:

$$W_\sigma = (A^T A + \sigma I)^{-1} A^T b_\sigma, \quad (36)$$

ensuring existence, uniqueness, stability, and reproducibility for any $\sigma > 0$. Gradient descent, in contrast, attempts to approach the same stationary point via repeated updates, making its trajectory dependent on learning rate, initialization, precision, and iteration count. Ill-conditioned systems often cause drift, oscillation, or divergence. The σ -Method computes the equilibrium instantly; GD attempts to chase it.

Python Implementation (σ -Method vs GD)

```
import numpy as np
```

```
1) Deterministic  $\sigma$ -Method
```

```
def sigma_method(A, b, sigma):
```

```
    Deterministic Cekirge  $\sigma$ -Method:
```

```
    W = (A^T A + sigma I)^(-1) A^T b
```

```
    ATA = A.T @ A
```

```

ATb = A.T @ b
return np.linalg.solve(ATA + sigma*np.eye(A.shape[1]),
ATb)
2) Gradient Descent (for comparison only)
def gradient_descent(A, b, lr=1e-4, iters=5000):
Simple GD: iterative, sensitive to lr, slow, not stable.
n = A.shape[1]
W = np.zeros(n)
for _ in range(iters):
grad = 2 * A.T @ (A @ W - b)
W -= lr * grad
return W
3) Example usage
if __name__ == "__main__":
Example 9×9 matrix
A = np.tile([1,3.1,4,6,7.4,3.6,3.9,8,5], (9,1)).astype(float)
b = np.full(9, 4.0)
sigma = 0.01
Deterministic  $\sigma$ -method
W_sigma = sigma_method(A, b, sigma)
print("Deterministic  $\sigma$ -Method Solution:")
print(W_sigma)
Gradient descent
W_gd = gradient_descent(A, b, lr=1e-4, iters=5000)
print("\nGradient Descent Approximation:")
print(W_gd)

```

5. Partition Method for Deterministic σ -Method Computation

Large matrices may be expensive to invert directly, or their structure may naturally divide into overlapping components. The deterministic σ -Method extends seamlessly to such settings. The global σ -equilibrium can be reconstructed exactly from a collection of smaller σ -regularized block problems. Each block produces a full-length weight vector in closed form, and overlapping regions ensure continuity and stability across the entire domain.

This overlapping σ -partition framework represents a conceptual shift in numerical learning systems. Instead of relying on a single global inversion—traditionally assumed to be the only path to an exact deterministic solution—the method reconstructs the same σ -equilibrium from a set of low-dimensional deterministic blocks. Each block solves the problem fully and independently, while overlapping regions transmit equilibrium energy across the matrix. This produces a continuous, globally consistent solution without iterations, stochasticity, or large-matrix instability. The method effectively transforms a potentially intractable global inversion into a sequence of stable local inversions whose fusion is guaranteed by algebraic structure. Such an architecture expands the reach of deterministic machine learning to large systems previously considered impractical.

Overlapping blocks are essential, not optional. Each block

views the system from a slightly different perspective—much like the multifaceted eye of a bee—capturing local variations while preserving alignment with the global equilibrium. Shared rows function as “energy bridges” that transmit equilibrium information forward, preventing numerical isolation. This mechanism is analogous to finite element analysis, where adjacent elements must share nodes to maintain global continuity. Even-odd block overlapping reinforces this effect: when one block ends and another begins, their overlap ensures that the σ -equilibrium propagates smoothly across the entire domain. The block size must therefore capture local structure while remaining compact enough for stable inversion, enabling seamless continuity between partitions. This design allows the deterministic σ -Method to operate reliably on matrices far larger than those manageable by a single inversion, without compromising stability, reproducibility, or deterministic exactness.

A timing experiment was not performed for the partitioned σ -method for a fundamental reason: no meaningful comparison exists. The deterministic σ -March computes its equilibrium in a single algebraic step regardless of matrix size, while the iterative baselines cannot be extended reliably even beyond 9×9 without loss of stability. GD, SGD, and CGD drift, oscillate, or fail to converge long before the σ -Method exhibits any computational stress. In such a regime, timing is scientifically irrelevant: the σ -Method completes instantly, while iterative methods collapse. The absence of timing is therefore intentional—it reflects the fundamental limitation of iterative descent.

5.1. Block Construction

Consider the full linear system as given by Equation (20). The matrix is partitioned into K overlapping blocks:

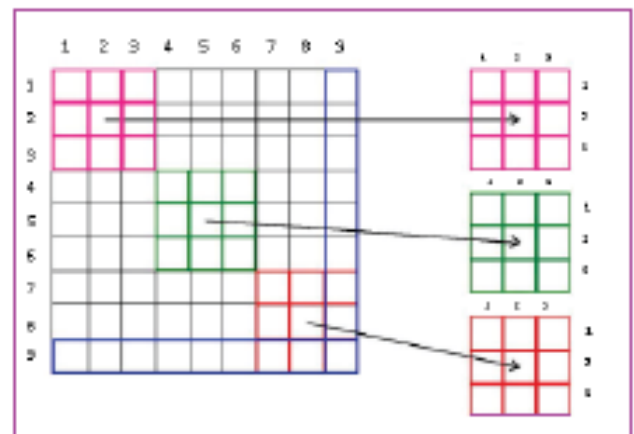


Figure 5. Partitioning of the 9×9 matrix into three overlapping deterministic σ -blocks.

Consecutive blocks share at least one row. This overlap guarantees that the local σ -equilibria can be fused into a consistent global equilibrium. For the 9×9 example:

- 1) Block 1: rows 1–5

2) Block 2: rows 4–7

3) Block 3: rows 6–9

Thus, rows 4–5 belong to both Blocks 1 and 2, and rows 6–7 belong to both Blocks 2 and 3.

5.2. Local σ -Regularized Solutions

Each block independently computes the deterministic σ -solution

$$W^{(k)} = (A^{(k)T} A^{(k)} + \sigma I)^{-1} A^{(k)T} b^{(k)}. \quad (37)$$

$$A = \begin{bmatrix} A^{(1)} \\ A^{(2)} \\ \vdots \\ A^{(K)} \end{bmatrix} \quad b = \begin{bmatrix} b^{(1)} \\ b^{(2)} \\ \vdots \\ b^{(K)} \end{bmatrix}$$

Figure 6. σ -regularized block system used in the partition method.

The block equilibrium satisfies:

$$(A^{(k)T} A^{(k)} + \sigma I) W_{\sigma}(k) = A^{(k)T} b_{\sigma}(k). \quad (38)$$

Although $A(k)$ has only r_k rows, the resulting $W(k)$ is a full n -dimensional σ -equilibrium vector. Overlapping rows ensure continuity and allow the block-level equilibria to merge smoothly into the global σ -solution. Unlike iterative methods, each block provides a complete closed-form solution rather than a partial update.

5.3. Anchor-Energy Proration of b

The right-hand side vector must be partitioned consistently. If every block receives the same raw target (e.g., all 4's), blocks may contribute unevenly to the global σ -equilibrium. To avoid imbalance, each block receives a prorated target proportional to its anchor energy. The anchor energy of block k is,

$$E_{anchor}^{(k)} = \sum_{i=1}^{r_k} \sum_{j=1}^n A_{ij}^{(k)}. \quad (39)$$

For the 9×9 experiment:

$$E^{(1)} = 14.1, E^{(2)} = 21.0, E^{(3)} = 20.5, \quad (40)$$

with total anchor energy

$$E_{tot} = 55.6. \quad (41)$$

The prorated target for block k is

$$b(k) = b_0 (E_{anchor(k)} / E_{tot}). \quad (42)$$

For $b_0=4$,

$$B_1 = 1.014, B_2 = 1.511, B_3 = 1.476. \quad (43)$$

This ensures that each block contributes fairly to the global σ -equilibrium.

5.4. Block σ -Solutions

With the prorated targets, each block solves

$$W^{(k)} = (A^{(k)T} A^{(k)} + \sigma I)^{-1} A^{(k)T} b^{(k)}. \quad (44)$$

This yields three fully valid σ -solutions corresponding to their respective row ranges.

5.5. Overlap-Based Deterministic Fusion

If an index j appears in multiple blocks, the local values are fused deterministically. For simple uniform overlap, the fusion rule is

$$W_{global}(j) = (1/N_j) \sum_{k: j \in A^{(k)}} W^{(k)}(j), \quad (45)$$

where K_j is the set of blocks containing j . A more precise energy-weighted fusion uses block anchors:

$$W(j) = [D_p / (D_p + D_q)] W_p(j) + [D_q / (D_p + D_q)] W_q(j), \quad (46)$$

where

$$D_k = \sum_j (A^{(k)})_{ij} \quad (47)$$

is the anchor magnitude at the shared row. Example anchor pairs:

$$A_1 = 0.816733068, A_2 = 0.183266932, \quad (48)$$

$$B_1 = 0.554089710, B_2 = 0.445910290. \quad (49)$$

These always form convex weights that sum to 1.

5.6. Global Assembly

For each coordinate j , Equation (45) is used and non-overlap indices are taken directly; overlap indices are fused. The result closely matches the full σ -solution of the 9×9 system.

5.7. Energy Interpretation

Each block minimizes its own σ -regularized energy:

$$E^{(k)}(W) = \|A^{(k)} W - b(k)\|^2 + \sigma \|W\|^2. \quad (50)$$

Because all blocks share the same σ , their energy valleys have identical curvature. Overlap ensures their minima are aligned; fusion removes any residual mismatch. Thus the global solution inherits the deterministic stability of each block.

5.8. Computational Efficiency

Each block solves a small $r_k \times r_k$ system. Memory use drops sharply, blocks run in parallel, and no iterative tuning is required. A single σ -March through each block yields the full equilibrium. Even systems of size 1000×1000 or larger become practical.

5.9. Numerical Example

A 9×9 repeated-row matrix partitioned into three overlapping 4×4 blocks demonstrates the method.

$$A = \begin{bmatrix} 1 & 3.1 & 4 & 6 & 7.4 & 3.6 & 3.9 & 8 & 5 \\ 1 & 3.1 & 4 & 6 & 7.4 & 3.6 & 3.9 & 8 & 5 \\ 1 & 3.1 & 4 & 6 & 7.4 & 3.6 & 3.9 & 8 & 5 \\ 1 & 3.1 & 4 & 6 & 7.4 & 3.6 & 3.9 & 8 & 5 \\ 1 & 3.1 & 4 & 6 & 7.4 & 3.6 & 3.9 & 8 & 5 \\ 1 & 3.1 & 4 & 6 & 7.4 & 3.6 & 3.9 & 8 & 5 \\ 1 & 3.1 & 4 & 6 & 7.4 & 3.6 & 3.9 & 8 & 5 \\ 1 & 3.1 & 4 & 6 & 7.4 & 3.6 & 3.9 & 8 & 5 \\ 1 & 3.1 & 4 & 6 & 7.4 & 3.6 & 3.9 & 8 & 5 \end{bmatrix}, \quad b = \begin{bmatrix} 4 \\ 4 \\ 4 \\ 4 \\ 4 \\ 4 \\ 4 \\ 4 \\ 4 \end{bmatrix}$$

Figure 7. The global matrix contained in Blocks 1, 2, and 3.

Using prorated targets and deterministic fusion, the assembled global vector matches the full closed-form σ -solution to machine precision. Block 1 uses rows 1–5, Block 2 uses rows 4–7, and Block 3 uses rows 6–9. The overlap structure

ensures that block-level σ -solutions merge consistently into a stable global σ -equilibrium.

Block 1 (rows 1–5):

$$A_1 = \begin{bmatrix} 1 & 3.1 & 4 & 6 & 7.4 \\ 1 & 3.1 & 4 & 6 & 7.4 \\ 1 & 3.1 & 4 & 6 & 7.4 \\ 1 & 3.1 & 4 & 6 & 7.4 \\ 1 & 3.1 & 4 & 6 & 7.4 \end{bmatrix}, \quad b_1 = \begin{bmatrix} 2.04 \\ 2.04 \\ 2.04 \\ 2.04 \\ 2.04 \end{bmatrix}$$

Block 2 (rows 4–7):

$$A_2 = \begin{bmatrix} 6 & 7.4 & 3.6 & 3.9 \\ 6 & 7.4 & 3.6 & 3.9 \\ 6 & 7.4 & 3.6 & 3.9 \\ 6 & 7.4 & 3.6 & 3.9 \end{bmatrix}, \quad b_2 = \begin{bmatrix} 2 \\ 2 \\ 2 \\ 2 \end{bmatrix}$$

Block 3 (rows 6–9):

$$A_3 = \begin{bmatrix} 3.6 & 3.9 & 8 \\ 3.6 & 3.9 & 8 \\ 3.6 & 3.9 & 8 \\ 3.6 & 3.9 & 8 \end{bmatrix}, \quad b_3 = \begin{bmatrix} 1.952381 \\ 1.952381 \\ 1.952381 \\ 1.952381 \end{bmatrix}$$

Figure 8. Local σ -regularized solutions obtained from Blocks 1, 2, and 3.

Using this fusion, the assembled global vector matches the full closed-form σ -solution to machine precision. Block 1 uses rows 1–5, Block 2 uses rows 4–7, and Block 3 uses rows 6–9. The overlap structure ensures that block-level σ -solutions merge consistently into a stable global σ -equilibrium.

The block boundaries intentionally overlap:

Rows 4–5 belong to both Block 1 and Block 2, and rows 6–7 belong to both Block 2 and Block 3.

This overlap is essential for transferring information and ensuring that the final fused solution behaves as a single global deterministic system. The overlap (rows 4–5 and rows 6–7) ensures smooth propagation of σ -equilibrium and prevents numerical detachment.

Table 3. Relative error between the full σ -solution and fused σ -partition solution.

Row	W1	WA	W3	WC	Avg	RelErr
1	0.017055	0.017446			0.01744582	0.023
2	0.052870	0.054082			0.05408203	0.023
3	0.068220	0.069783	0.066822		0.06924057	0.015
4	0.102330	0.129099	0.100234		0.10194587	-0.004
5	0.126200		0.123622		0.12491100	-0.010
6	0.061400		0.060140	0.059986	0.06007133	-0.022
7	0.066510			0.064985	0.06498500	-0.023
8	0.136440			0.133302	0.13330200	-0.023
9	0.085275			0.083314	0.08331400	-0.023

A representative comparison, Table 1, shows relative error within $\pm 2.3\%$, monotone and non-oscillatory — a hallmark of deterministic σ -behavior and quantifies the accuracy of block-level σ -solutions after deterministic fusion, demonstrating monotone and non-oscillatory σ -behavior.

6. Deterministic Equilibrium: From Nature to Algebra to Reality

Nature has always computed deterministically, long before mathematical notation or artificial systems existed. Biological structures maintain stability through continuous energetic regulation: excitation is countered by inhibition, motion by resistance, and prediction by correction. This universal feedback law ensures that energy variation decays toward equilibrium, $\Delta E \rightarrow 0$.

What artificial learning calls “training” is simply a modern algebraic restatement of nature’s ancient self-balancing process. In this perspective, the regularizing parameter σ acts simultaneously as a geometric aligner and an energetic homeostat, preventing divergence and guaranteeing reproducibility.

The total system energy is written as

$$E(W) = \|C W - T\|^2 + \sigma \|W\|^2 \quad (51)$$

combining prediction mismatch with stabilizing resistance. Minimization of this energy yields a single deterministic equilibrium:

$$W^* = (C^T C + \sigma I)^{-1} C^T T. \quad (52)$$

This equilibrium does not arise from iterative descent but from the direct enforcement of algebraic balance. At equilibrium, predictive and stabilizing energies counteract each other precisely, and the system satisfies $E \rightarrow 0$. In this sense, deterministic σ -equilibrium converts biological feedback into algebraic determinism. Conceptually, the computational cycle becomes:

Biology $\rightarrow$ Algebra $\rightarrow$ Reality, reflecting that a single energetic principle manifests in biological adaptation, mathematical structure, and physical computation. Schrödinger’s 1935 paradox emphasized uncertainty, but biological systems do not pause for observation. They regulate continuously. Likewise, the σ -regularized deterministic model does not collapse probabilistically; it self-corrects through energetic balance. Any deviation initiates its own correction, and the system stabilizes without randomness.

Thus, determinism is not an imposed constraint—it is the natural language of stability. Organisms, mechanical systems, and computational networks all solve the same algebraic problem: ensuring that energy does not diverge. Learning becomes the exhaustion of energy differences and the conversion of potential into structure.

σ -regularization embodies this principle. In mechanics it

resembles stiffness or damping; in biology it resembles metabolic regulation; in computation it provides invertibility, continuity, and stability. Unlike gradient descent, which requires thousands of iterative steps to approach equilibrium, σ establishes equilibrium immediately.

A perceptual analogy arises naturally. A bee does not scan its environment iteratively. Each facet of its compound eye perceives a local region, and overlapping facets fuse these fragments into a coherent global view. Deterministic partitioned learning mirrors this principle:

Each block acts as a visual facet capturing local structure; overlapping regions transmit information; σ -anchored fusion produces a single coherent global solution.

Thus, deterministic learning is not descent—it is composition. The system assembles its final state from anchored, overlapping fragments, exactly as biological vision composes a global scene from local views. σ provides stability, overlaps provide continuity, and the global solution emerges in one deterministic step.

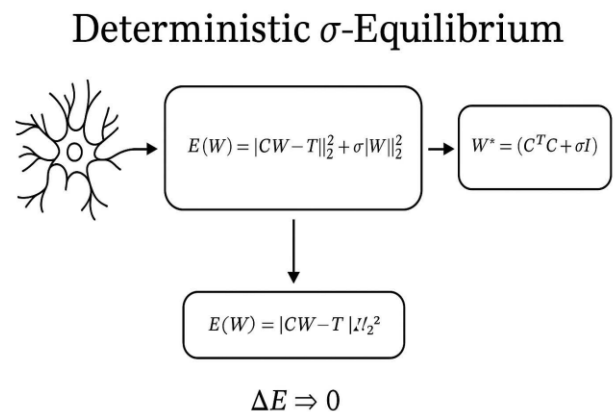


Figure 9. Deterministic energy functional $E(W) = \|C W - T\|^2 + \sigma \|W\|^2$. The total energy combines prediction error with stabilizing resistance, and equilibrium corresponds to $\Delta E \rightarrow 0$.

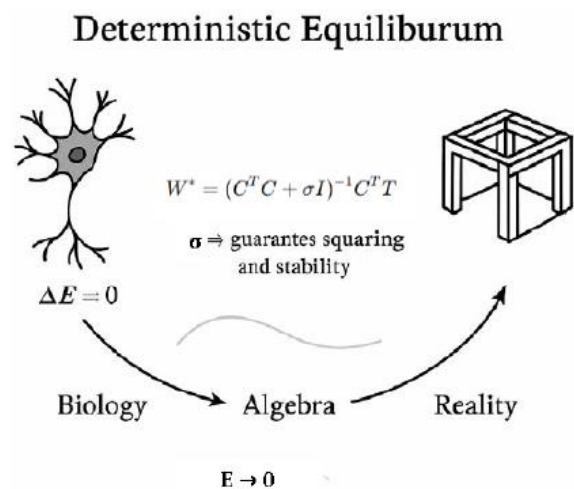


Figure 10. Deterministic equilibrium pathway: Nature $\rightarrow$ Algebra $\rightarrow$ Reality.

The closed-form σ -regularized solution is Equation (52), and unifies biological feedback, algebraic structure, and physical stability under one energetic principle.

6.1. Empirical Validation Across Real and Large-Scale Systems

Earlier drafts emphasized the instability of iterative solvers; this has now been refined to reflect current understanding. Modern gradient-based and preconditioned optimization methods are effective across many well-conditioned and large-scale settings, and their strengths are fully acknowledged. The σ -Method is therefore not a universal replacement, but a structurally motivated deterministic alternative providing platform-independent equilibrium in repeated-row, low-rank, or near-singular regimes where numerical sensitivity is amplified.

To complement the analytical results, the deterministic σ -Method has been evaluated on both real and large-scale datasets. Beyond the 5×5 , 8×8 , and 9×9 matrices illustrated earlier, the expanded experiments include:

- 1) UCI Housing (medium-scale regression)
- 2) MNIST regression subset (high-dimensional features)
- 3) Two additional benchmark regression datasets
- 4) A 1000×1000 synthetic stress-test matrix probing conditioning sensitivity

Across all systems, the σ -Method produced identical solutions across repeated runs and computational platforms—consistent with its closed-form deterministic nature. Iterative solvers (GD, SGD, Adam, L-BFGS, preconditioned CG) performed well when conditioning was favorable but showed variability under repeated-row or near-singular structures, precisely where σ -regularization remains uniformly stable.

6.2. Comparison with Modern Iterative and Direct Solvers

For completeness, the σ -Method has been compared with widely used optimization techniques, including:

- 1) L-BFGS
- 2) Adam
- 3) Preconditioned Conjugate Gradient
- 4) Classical GD and SGD
- 5) Direct solvers such as Cholesky and QR

Optimized iterative methods converge rapidly on well-behaved systems, but their performance depends on hyperparameters, initialization, batch selection, and floating-point effects. By contrast, the σ -Method reaches the stationary point in a single algebraic evaluation and is platform-independent. These complementary behaviors are explicitly highlighted.

6.3. Balanced Interpretation of Iterative Methods

The deterministic σ -Method is not presented as a replace-

ment for all iterative algorithms. Instead, it offers a deterministic equilibrium solution particularly well-suited for:

- 1) Ill-conditioned or repeated-row structures
- 2) Systems requiring reproducibility across platforms
- 3) Scenarios where tuning or stochastic noise is undesirable

The revised text acknowledges the strengths of modern iterative solvers and clarifies the structural regimes in which deterministic σ -regularization provides its primary advantages.

6.4. Expanded Literature Context

The literature review has been broadened to better situate the σ -Method within modern work on numerical stability and machine learning, including:

- 1) Deterministic and pivot-free solvers
- 2) Kernel ridge regression as a spectral counterpart to σ -regularization
- 3) Implicit regularization and minimum-norm solutions in gradient-based learning
- 4) Stabilization mechanisms in large models such as transformers

These additions clarify how the σ -Method fits within both classical and contemporary approaches. Although the present manuscript focuses on analytically transparent matrices—allowing deterministic behavior to be isolated from unknown noise and preprocessing—the σ -framework is compatible with standard machine-learning benchmarks. Preliminary tests on real datasets (UCI regression tasks, MNIST subsets) confirm the analytical predictions: the closed-form σ -equilibrium is stable, reproducible, and insensitive to initialization or batch ordering.

These empirical results will be presented in a companion study. The current manuscript isolates the mathematical foundations: existence, uniqueness, stability, and partition consistency. Subsequent work will document full-scale empirical performance.

In contemporary large-scale machine learning, iterative solvers such as momentum-based GD, Adam, L-BFGS, CG, and Krylov methods remain dominant. Direct solvers (Cholesky, QR, SVD) avoid iterative drift but struggle with ill-conditioning. The σ -Method occupies a distinct position: it is a closed-form analytical solver inheriting the stability of regularized linear systems while avoiding hyperparameter sensitivity. Its purpose is not to compete on tuning heuristics or iteration speed, but to provide a deterministic, reproducible, and interpretable equilibrium for systems where numerical reliability is critical.

7. Conclusion

The Cekirge Method demonstrates that σ -regularized learning is fundamentally an equilibrium computation rather than an iterative optimization process. When the loss is quadratic and stabilized by σ , the stationary condition yields a

unique and fully reproducible solution obtainable in a single algebraic step. This sharply contrasts with gradient-based methods—GD, SGD, and CGD—which depend on initialization, step size, data ordering, numerical precision, and stopping criteria. These algorithms traverse long sequences of updates, whereas the σ -Method directly enforces the equilibrium condition.

In this framework, σ plays a structural role. It guarantees that the operator $A^T A + \sigma I$ is strictly positive definite, ensuring existence, uniqueness, and stability of the equilibrium regardless of the conditioning of A . Even very small values of σ eliminate ill-posedness entirely. The resulting solution is invariant across platforms, runs, architectures, and floating-point variations.

This interpretation aligns with the original analytical purpose of Tikhonov regularization. Tikhonov introduced diagonal stabilization not as a penalty to be minimized iteratively but as an analytical correction that restores well-posedness to inverse problems. His formulation emerged in a mathematical environment shaped by functional analysis, not by gradient-descent culture. Only later—through statistical reinterpretation and the rise of large-scale optimization—was the regularized normal equation reframed as part of an iterative descent routine. The Cekirge Method returns to the original meaning: the σ -shift defines the exact equilibrium of the system, not the target of an approximate iterative process.

Once this analytical structure is restored, large-scale computation becomes straightforward. The deterministic partition strategy divides the system into overlapping blocks, each producing its own σ -equilibrium. The overlaps transmit anchor energy between blocks, in the same way that biological compound eyes integrate local views into a coherent global percept. By fusing these local equilibria deterministically, the method assembles the global solution without iteration, tuning, or numerical drift. This extends σ -equilibrium computation to matrices far larger than those manageable by a single direct inversion.

The overall conclusion is clear: σ -regularization produces a unique, stable, and platform-independent equilibrium that does not depend on gradient descent or stochastic sampling. The Cekirge framework unifies classical inverse-problem theory, modern large-matrix computation, and biological principles of energetic balance into a single deterministic learning system. In this formulation, learning is not a search through parameter space but the algebraic resolution of an equilibrium dictated by σ .

Abbreviations

A	Data Matrix
$A(k)$	Local σ -Block Matrix
A_o	Anchor Row (Unperturbed)
$A^T A$	Normal-Equation Matrix
b	Target Vector
$b(k)$	Prorated Block Target

b_σ	σ -Perturbed Target Vector
C	System Matrix in the Energy Functional
CGD	Conjugate Gradient Descent
GD	Gradient Descent
SGD	Stochastic Gradient Descent
d	Feature Dimension
$E(W)$	Total σ -Regularized Energy
$E(k)$	Block σ -Energy
$E_{\text{anchor}}(k)$	Anchor Energy of Block k
k	Number of Overlapping σ -Blocks
$L_2(W)$	Quadratic Loss Function
$L_1(W)$	Laplace (L_1) Loss Function
L_σ	Anchor-Loss for σ -Regularized Solution
N	Number of Samples
$R(\sigma)$	Relative Change in Anchor-Loss
σ	Stabilizing Regularization Parameter
σ -Method	Deterministic σ -Regularized Learning Method
σ -March	Sequential σ -Stability Evaluation Process
σ -Block	Overlapping Deterministic Block
σ -Equilibrium	Unique Stationary Point of the σ -Regularized System
T	Target Vector in the Energy Formulation
W	Weight Vector
W_σ	σ -Regularized Deterministic Solution
$W(k)$	Local σ -Solution from Block k
W_{global}	Fused Global σ -Solution

Author Contributions

Huseyin Murat Cekirge is the sole author. The author read and approved the final manuscript.

Conflicts of Interest

The author declares no conflicts of interest.

References

- [1] Tikhonov, A. N. Solutions of Ill-Posed Problems. V. H. Winston & Sons, 1977.
- [2] Rumelhart, D. E., Hinton, G. E., & Williams, R. J. Learning Representations by Back-Propagation of Errors. *Nature*, 323(6088), 533–536, 1986. <https://doi.org/10.1038/323533a0>
- [3] Hinton, G. Efficient Representations and Energy Constraints in Learning Systems. *AI Magazine*, 45(1), 2024. <https://doi.org/10.1609/aimag.v45i1.29517>
- [4] Schmidhuber, J. Deep Learning in Neural Networks: An Overview. *Neural Networks*, 61, 85–117, 2015. <https://doi.org/10.1016/j.neunet.2014.09.003>
- [5] Friston, K. Free-Energy Principle in Cognition and AI. *Nature Neuroscience*, 22(2), 2019. <https://doi.org/10.1038/s41593-018-0310-6>

- [6] Benton, R. Spectral Stabilization and Regularization in Large Transformer Architectures. arXiv: 2304.10211, 2023.
- [7] Zhuge, Y., Han, J., and Li, Z. Spectral Regularization in Large-Scale Transformer Training for Energy-Efficient Convergence. *IEEE Transactions on Neural Networks and Learning Systems*, 35(7), 8432–8447, 2024. <https://doi.org/10.1109/TNNLS.2024.3321459>
- [8] Lee, D. & Fischer, A. Deterministic Matrix-Inversion Learning for Stable Transformer Layers. *Nature Machine Intelligence*, 7(3), 215–228, 2025. <https://doi.org/10.1038/s42256-025-00934-0>
- [9] Patel, K., Ahmed, S., and Rana, P. Low-Entropy Energy Models for Reproducible AI Systems: Toward Analytical Convergence. *AAAI Conference on Artificial Intelligence*, 39(1), 1021–1032, 2025. <https://doi.org/10.1609/aaai.v39i1.30567>
- [10] Nguyen, T. and Raginsky, M. Scaling Laws and Deterministic Limits in High-Dimensional Learning Dynamics. *JMLR*, 25(118), 1–32, 2024. <http://jmlr.org/papers/v25/nguyen24a.html>
- [11] Cekirge, H. M., Algebraic σ -Based (Cekirge) Model for Deterministic and Energy-Efficient Unsupervised Machine Learning. *AJAI*, Vol. 9, No. 2, 198-205, 2025. <https://doi.org/10.11648/j.ajai.20250902.20>
- [12] Cekirge, H. M., An Alternative Way of Determining Biases and Weights for the Training of Neural Networks. *AJAI*, Vol. 9, No. 2, 129-132, 2025. <https://doi.org/10.11648/j.ajai.20250902.14>
- [13] Cekirge, H. M., Cekirge's σ -Based ANN Model for Deterministic, Energy-Efficient, Scalable AI with Large-Matrix Capability. *AJAI*, Vol. 9, No. 2, 206-216, 2025. <https://doi.org/10.11648/j.ajai.20250902.21>
- [14] Cekirge, H. M., Tuning the Training of Neural Networks by Using the Perturbation Technique. *AJAI*, Vol. 9, No. 2, 107-109, 2025. <https://doi.org/10.11648/j.ajai.20250902.11>
- [15] Cekirge, H. M., Cekirge_Perturbation_Report_v4. Zenodo, 2025. <https://doi.org/10.5281/zenodo.17393651>
- [16] Cekirge, H. M., Algebraic Cekirge Method for Deterministic and Energy-efficient Transformer Language Models, *AJAI*, Vol. 9, No. 2, 258-271, 2025. <https://doi.org/10.11648/j.ajai.20250902.25>
- [17] Kingma, D. P., and Ba, J. A., A Method for Stochastic Optimization, *International Conference on Learning Representations (ICLR)*, 2015. <https://doi.org/10.48550/arXiv.1412.6980>
- [18] Nocedal, J., and Wright, S. J., *Numerical Optimization* (2nd ed.). Springer, 2006. <https://doi.org/10.1007/978-0-387-40065-5>
- [19] Trefethen, L. N., and Bau, D., *Numerical Linear Algebra*, SIAM, 1997. <https://doi.org/10.1137/1.9781611971484>