

Article

Not peer-reviewed version

CLSP: Linear Algebra Foundations of a Modular Two-Step Convex Optimization-Based Estimator for Ill-Posed Problems

[Ilya Bolotov](#) *

Posted Date: 26 September 2025

doi: 10.20944/preprints202509.2192.v1

Keywords: convex optimization; modular estimators; least-squares; generalized inverse; regularization; normalized RMSE; Monte Carlo simulation



Preprints.org is a free multidisciplinary platform providing preprint service that is dedicated to making early versions of research outputs permanently available and citable. Preprints posted at Preprints.org appear in Web of Science, Crossref, Google Scholar, Scilit, Europe PMC.

Copyright: This open access article is published under a Creative Commons CC BY 4.0 license, which permit the free download, distribution, and reuse, provided that the author and preprint are cited in any reuse.

Disclaimer/Publisher's Note: The statements, opinions, and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions, or products referred to in the content.

Article

CLSP: Linear Algebra Foundations of a Modular Two-Step Convex Optimization-Based Estimator for Ill-Posed Problems

Ilya Bolotov 

Prague University of Economics and Business, Faculty of International Relations, W. Churchill Sq. 1938/4, 130 67 Prague 3 – Žižkov, Czech Republic; ilya.bolotov@vse.cz or ekonobol@gmail.com

Abstract

This paper develops the linear-algebraic foundations of the Convex Least Squares Programming (CLSP) estimator and constructs its modular two-step convex optimization framework, capable of addressing ill-posed and underdetermined problems. After reformulating a problem in its canonical form, $\mathbf{A}^{(r)}\mathbf{z}^{(r)} = \mathbf{b}$, Step 1 yields an iterated (if $r > 1$) minimum-norm least-squares estimate $\hat{\mathbf{z}}^{(r)} = (\mathbf{A}_Z^{(r)})^+ \mathbf{b}$ on a constrained subspace defined by a symmetric idempotent $\mathbf{Z}$ (reducing to the Moore-Penrose pseudoinverse when $\mathbf{Z} = \mathbf{I}$). The optional Step 2 corrects $\hat{\mathbf{z}}^{(r)}$ by solving a convex program, which penalizes deviations using a Lasso/Ridge/Elastic net-based scheme parameterized by $\alpha \in [0, 1]$ and yields $\hat{\mathbf{z}}^*$. The second step guarantees a unique solution for $\alpha \in (0, 1]$ and coincides with the Minimum-Norm BLUE (MNBLUE) when $\alpha = 1$. This paper also proposes an analysis of numerical stability and CLSP-specific goodness-of-fit statistics, such as partial R^2 , normalized RMSE (NRMSE), Monte Carlo t -tests for the mean of NRMSE, and condition-number-based confidence bands. The three special CLSP problem cases are then tested in a 50,000-iteration Monte Carlo experiment and on simulated numerical examples. The estimator has a wide range of applications, including interpolating input-output tables and structural matrices.

Keywords: convex optimization; modular estimators; least-squares; generalized inverse; regularization; normalized RMSE; Monte Carlo simulation

MSC: 90C25, 65F22, 65F20, 65F35, 65C05

JEL Classification: C13, C61, C63, C15

1. Introduction

As existing literature on optimization — for example, Nocedal and Wright [1], pp. 1–9 and Boyd and Vandenberghe [2], pp. 1–2 — points out, constrained optimization plays an important role not only in natural sciences, where it was first developed by Euler, Fermat, Lagrange, Laplace, Gauss, Cauchy, Weierstrass, and others, but also in economics and its numerous applications, such as econometrics and operations research, including managerial planning at different levels, logistics optimization, and resource allocation. The history of modern constrained optimization began in the late 1940s to 1950s, when economics, already heavily reliant on calculus since the marginalist revolution of 1871–1874 [3], pp. 43–168, and the emerging discipline of econometrics, standing on the shoulders of model fitting [4] (pp. 438–457), internalized a new array of quantitative methods referred to as programming after the U.S. Army's logistic jargon. Programming — including classical, linear, quadratic, and dynamic variants — originated with Kantorovich in economic planning in the USSR and was independently discovered and developed into a mathematical discipline by Dantzig in the United States (who introduced the bulk of terminology and the simplex algorithm) [5] [6], pp. 1–51. Since the pivotal proceedings published by Koopmans [7], programming became an integral

part of mathematical economics [8–10], in both microeconomic analysis (an overview is provided in Intriligator and Arrow [11], pp. 76–91) and macroeconomic modeling, such as Leontief’s input-output framework outlined in Koopmans [7], pp. 132–173 and, in an extended form, in Dorfman et al. [12], pp. 204–264.

Calculus, model fitting, and programming are, however, subject to strong assumptions, including function smoothness, full matrix ranks, well-posedness, well-conditionedness, feasibility, non-degeneracy, and non-cycling (in the case of the simplex algorithm), as summarized by Nocedal and Wright [1], pp. 304–354, 381–382 and Intriligator and Arrow [11], pp. 15–76. These limitations have been partially circumvented by developing more strategy-based than analysis-based, computationally intensive constraint programming algorithms; consult Frühwirth and Abdennadher [13] and Rossi et al. [14] for an overview. Still, specific cases of constrained optimization problems in econometric applications, such as incorporating the textbook two-quarter business-cycle criterion into a general linear model [15] or solving underdetermined linear programs to estimate an investment matrix [16], remain unresolved within the existing methodologies and algorithmic frameworks. In response, this text develops a modular two-step convex programming methodology to formalize such problems as a linear system of the form $\mathbf{A}\mathbf{z} = \mathbf{b}$, including linear(ized) constraints $\mathbf{C}\mathbf{x} + \mathbf{S}\mathbf{y}^* = \mathbf{b}_{1:k}$, with $k \leq m$, where $\mathbf{A} \in \mathbb{R}^{m \times n}$ is a partitioned matrix, $\mathbf{C} \in \mathbb{R}^{k \times p}$ and $\mathbf{S} \in \mathbb{R}^{k \times (n-p)}$ are its submatrices, $\mathbf{b} \in \mathbb{R}^m$ is a known vector, and $\mathbf{z} \in \mathbb{R}^n$ is the vector of unknowns, which includes both the target solution subvector $\mathbf{x} \in \mathbb{R}^p$ and an auxiliary slack-surplus subvector $\mathbf{y}^* \in \mathbb{R}^{n-p}$, introduced to accommodate inequalities in constraints. The focus is on its linear algebra foundations with corresponding theorems and lemmas, followed by a discussion of numerical stability, goodness of fit, and potential applications.

The text is organized as follows. Section 2 summarizes the historical development and recent advances in convex optimization, motivating the formulation of a new methodology. Section 3 presents the formalization of the estimator, while Section 4 develops a sensitivity analysis for $\mathbf{C}$. Section 5 introduces goodness-of-fit statistics, including the normalized root mean square error (NRMSE) and its sample-mean test. Section 6 presents special cases, and Section 7 reports the results of a Monte Carlo experiment and solves simulated problems via CLSP-based Python 3 modules. Section 8 concludes with a general discussion. Throughout, the following notation is used: bold uppercase letters (e.g., $\mathbf{A}, \mathbf{B}, \mathbf{C}, \mathbf{G}, \mathbf{M}, \mathbf{S}$) denote matrices; bold lowercase letters (e.g., $\mathbf{b}, \mathbf{x}, \mathbf{y}, \mathbf{z}$) denote column vectors; italic letters (e.g., $i, j, k, m, n, p, \lambda, \omega$) denote scalars; the transpose of a matrix is denoted by the superscript $\top$ (e.g., $\mathbf{X}^\top$); the inverse by -1 (e.g., $\mathbf{Q}^{-1}$); generalized inverses are indexed using curly braces (e.g., $\mathbf{G}^{\{1,2\}}$); the Moore-Penrose pseudoinverse is denoted by dagger (e.g., $\mathbf{A}^\dagger$); ℓ_p norms, where $1 \leq p \leq \infty$, are denoted by double bars (e.g., $\|\cdot\|_2$); and condition numbers by $\kappa(\cdot)$. All functions, scalars, vectors, and matrices are defined over the real numbers ($\mathbb{R}$).

2. Historical and Conceptual Background of the CLSP Framework

Convex optimization — formally, $\mathbf{x}^* = \arg \min_{\mathbf{x} \in \mathbb{R}^n} f(\mathbf{x})$, subject to convex inequality constraints $g_i(\mathbf{x}) \leq 0$ and affine equality constraints $h_j(\mathbf{x}) = 0$ [2], pp. 7, 127–129, encompassing, among others, least squares, linear programming, and quadratic programming — has its roots in *differential calculus*, a concept developed independently by Newton — who was influenced by Wallis and Barrow — in 1665–1666 (published posthumously in 1736), and by Leibniz, who drew among others on Pascal, in 1673–1675 (published in 1684 and 1686) [17], and, in its constrained form, by Lagrange in the 1760s. Calculus-based optimization dominated analytical methods in mathematics and physics until the 19th century, during which period it was superseded in specific applications by model-fitting techniques, driven by empirical inconsistencies in classical mechanics, particularly in celestial motion [18], pp. 11–61. Subsequent contributions by Euler, Laplace, Cauchy, and Weierstrass helped shape the modern algorithmic treatment of calculus-based optimization, as presented in, for example, Nocedal and Wright [1], pp. 10–29 and, for economic applications, in Intriligator and Arrow [11], pp. 53–72 and Ok [19], pp. 670–745. In (micro)economic theory, Lagrangian optimization was introduced by Edgeworth, who applied it in 1877 to utility maximization, in 1879 to utilitarian welfare distribution, and in 1881

to general equilibrium via the contract curve [20]. Since the 19th century, optimizing differentiable scalar-valued functions has become the foundation for the fitting techniques in mathematical statistics [4], pp. 494–510.

Stigler [18], pp. 11–61 traced the pre-19th century history leading to the invention of the *least squares method* (named after the original French term *la méthode des moindres carrés*, henceforth abbreviated as LS) by Legendre in 1805, Adrain in 1808–1809, and independently by Gauss in 1809 (although the latter claimed to have been using the method since 1795; a rigorous discussion of the priority dispute between Gauss and Legendre, a well-known case study in the history of scientific priority, is provided in Plackett [21] and Stigler [22] [23], pp. 320–331). The method's formalism evolved from Gauss's, Yule's, and Laplace's pre-matrix notations to Pearson's, Fisher's, Frisch's, Aitken's (who is believed to anticipate the Gauss-Markov theorem [24]), and Bartlett's use of Cayley's 1850s and subsequent methodologies, with the modern $\mathbf{y} = \mathbf{X}\mathbf{b}$ notation, where $\mathbf{X}$ is a matrix and $\mathbf{y}, \mathbf{b}$ are vectors, introduced around 1950 and already in use by Durbin and Watson, among others [25]. LS became an essential tool at the dawn of modern econometrics and its application in linear regression (ordinary least squares, OLS) intensified with the creation of the Cowles Commission in the 1930s–1940s U.S. and the works of Fisher, Koopmans, Haavelmo, Stone, Prest, and others [26], pp. 117–145, [27], pp. 7–12. Contrary to the stricter assumptions of pure calculus-based optimization, LS (a) enabled modeling over discrete or irregular data, and (b) incorporated residuals, moving beyond exact curve fitting, while still ensuring minimum variance among unbiased linear estimators under the Gauss–Markov theorem.

Early efforts to estimate the normal equations $\mathbf{X}^\top \mathbf{y} = \mathbf{X}^\top \mathbf{X} \mathbf{b}$, where $\mathbf{X}^\top \mathbf{X}$ is the sum-of-squares and cross-products matrix, spurred the development of generalized inverses, numerically stable matrix decompositions, and regularization techniques to tackle *rank-deficient* and *ill-posed problems*; consult Gentle [4], pp. 126–154, 217–301 for a comprehensive modern overview. Their interconnection — namely, for the Moore–Penrose pseudoinverse $\mathbf{A}^\dagger$, which has been computed predominantly through singular value decomposition (SVD) since the 1960s, as introduced by Golub and Kahan [28] (a numerical algorithm for the weighted pseudoinverse was provided in the 1980s by Eldén [29]) and was demonstrated, among others, in the 1970s by Ward [30], p. 34 and, more profoundly, in the 2010s by Barata and Hussein [31], pp. 152–154 to be the limit of a Tikhonov regularization formula $\mathbf{A}^\dagger = \lim_{\mu \rightarrow 0} (\mathbf{A}^\top \mathbf{A} + \mu \mathbf{I})^{-1} \mathbf{A}^\top$, where $\mathbf{I} \in \mathbb{R}^{n \times n}$ is an identity matrix for $\mathbf{A} \in \mathbb{R}^{m \times n}$ and μ is a parameter (for a detailed presentation of the regularization in its original notation, consult Tikhonov et al. [32], pp. 10–67) — has become a standard perspective in modern numerical analysis and forms the foundation of the CLSP framework, the first step of which is $\hat{\mathbf{z}} = \mathbf{G}\mathbf{b}$, where $\mathbf{G}$ is either $\mathbf{A}^\dagger \in \mathbb{R}^{n \times m}$ or a constrained reflexive inverse $\mathbf{G}^{\{1,2\}} \in \mathbb{R}^{n \times m}$.

As documented in the seminal works of Rao and Mitra [33], pp. vii–viii and Ben-Israel and Greville [34], pp. 4–5, 370–374, the first pseudoinverses (original term, now typically reserved for $\mathbf{A}^\dagger$) or generalized inverses (current general term) emerged in the theory of integral operators — introduced by Fredholm in 1903 and further developed by Hurwitz in 1912 — and, implicitly, in the theory of differential operators by Hilbert in 1904, whose work was subsequently extended by Myller in 1906, Westfall in 1909, Bounitzky in 1909, Elliott in 1928, and Reid in 1931. The cited authors attribute their first application to matrices to Moore in 1920 under the term *general reciprocals* [35] (though some sources suggest he may have formulated the idea as early as 1906), with independent formulations by Siegel in 1937 and Rao in 1955, and generalizations to singular operators by Tseng in 1933, 1949, and 1956, Murray and von Neumann in 1936, and Atkinson in 1952–1953. A theoretical consolidation occurred with Bjerhammar's [36] rediscovery of Moore's formulation, followed by Penrose's [37,38] introduction of four conditions defining the unique least-squares minimum-norm generalized inverse $\mathbf{G} \equiv \mathbf{A}^\dagger$, which, for $\mathbf{G} \in \mathbb{R}^{n \times m}$, can be expressed in the following way:

$$(i) \mathbf{AGA} = \mathbf{A} \quad (ii) \mathbf{GAG} = \mathbf{G} \quad (iii) (\mathbf{AG})^\top = \mathbf{AG} \quad (iv) (\mathbf{GA})^\top = \mathbf{GA} \quad (1)$$

In econometrics, it led to conditionally unbiased (minimum bias) estimators [39] and a redefinition of the Gauss-Markov theorem [40] in the 1960s, superseding earlier efforts [41].

Further contributions to the theory, calculation, and diverse applications of $\mathbf{G}$ were made in the 1950s–early 1970s by Rao and Mitra [33,42] (a synthesis of the authors’ previous works from 1955–1969), Golub and Kahan [28], Ben-Israel and Greville [34], Greville [43,44,45,46], Cline [47], Cline and Greville [48], Ben-Israel and Cohen [49], and Lewis and Newman [50] (the most cited sources on the topic in the Web of Science database, although the list is not exhaustive; consult Rao and Mitra [33], pp. vii–viii for an extended one), which, among others, introduced the concept of a $\{\cdot\}$ -generalized inverse $\mathbf{G} \in \mathbf{A}^{\{\cdot\}}$, where $\mathbf{A}^{\{\cdot\}} \subseteq \mathbb{R}^{n \times m}$ is the space of the $\{\cdot\}$ -generalized inverses of $\mathbf{A} \in \mathbb{R}^{m \times n}$, — hereafter, $\{i, j\}$ -inverses for $i > 1$, $\{i, j, k\}$ -inverses, Drazin inverses, and higher-order $\mathbf{G}$ are disregarded because of their inapplicability to the CLSP framework — by satisfying a subset or all of the (i)–(iv) conditions in Equation (1), as described in Table 1. Formally, for $\mathbf{G}^{\{1\}} \in \mathbf{A}^{\{1\}}$, $\mathbf{I}_n \in \mathbb{R}^{n \times n}$, $\mathbf{I}_m \in \mathbb{R}^{m \times m}$, and arbitrary matrices $\mathbf{U}, \mathbf{V} \in \mathbb{R}^{n \times m}$, any $\{1\}$ -inverse, including all $\{1, j\}$ -inverses and $\mathbf{A}^+$, can be expressed as $\mathbf{G} = \mathbf{G}^{\{1\}} + (\mathbf{I}_n - \mathbf{G}^{\{1\}}\mathbf{A})\mathbf{U} + \mathbf{V}(\mathbf{I}_m - \mathbf{A}\mathbf{G}^{\{1\}})$, out of which only $\mathbf{G} \equiv \mathbf{A}^+$ and a constrained $\mathbf{G} \equiv \mathbf{G}^{\{1,2\}} \in \mathbb{R}^{n \times m}$ — here, the Bott-Duffin inverse [51], expressed in modern notation as $\mathbf{A}_S^+ = \mathbf{P}_S^L(\mathbf{P}_S^L\mathbf{A}\mathbf{P}_S^R)^+ \mathbf{P}_S^R$, where $\mathbf{P}_S^L \in \mathbb{R}^{m \times m}$ and $\mathbf{P}_S^R \in \mathbb{R}^{n \times n}$ are orthogonal (perpendicular) projectors onto the rows and columns defined by an orthonormal matrix $\mathbf{S}$ — can be uniquely defined ($\mathbf{A}_S^+$ under certain conditions) and qualify as minimum-norm unbiased estimators in the sense of Chipman [39] and Price [40].

Table 1. Selected types of generalized inverses $\mathbf{G} \in \mathbf{A}^{\{\cdot\}}$, where $\mathbf{A}^{\{\cdot\}} \subseteq \mathbb{R}^{n \times m}$, $\mathbf{x} \in \mathbb{R}^n$, and $\mathbf{b} \in \mathbb{R}^m$.

Type of Inverse	Terminology	Properties
$\{1\}$	Equation Solving Inverse	$\mathbf{G} \in \mathbf{A}^{\{1\}}$ iff $\mathbf{AGA} = \mathbf{A}$. Hence, $\mathbf{Gb}$ is a solution to $\mathbf{Ax} = \mathbf{b}$ for all $\mathbf{b} \in \mathcal{R}(\mathbf{A})$.
$\{1, 2\}$	Reflexive Inverse	$\mathbf{G} \in \mathbf{A}^{\{1,2\}}$ iff $\mathbf{AGA} = \mathbf{A}$ and $\mathbf{GAG} = \mathbf{G}$. Each $\{1, 2\}$ -inverse defines a direct sum $\mathcal{N}(\mathbf{A}) \oplus \mathcal{R}(\mathbf{G}) = \mathbb{R}^n$, $\mathcal{R}(\mathbf{A}) \oplus \mathcal{N}(\mathbf{G}) = \mathbb{R}^m$. For complementary subspaces $(\mathcal{N}, \mathcal{R})$, $\mathbf{G}_{\mathcal{N}, \mathcal{R}}$ is unique, with $\mathcal{R}(\mathbf{G}_{\mathcal{N}, \mathcal{R}}) = \mathcal{R}$ and $\mathcal{N}(\mathbf{G}_{\mathcal{N}, \mathcal{R}}) = \mathcal{N}$.
$\{1, 3\}$	Least Squares Inverse	$\mathbf{G} \in \mathbf{A}^{\{1,3\}}$ iff $\mathbf{AGA} = \mathbf{A}$ and $(\mathbf{AG})^\top = \mathbf{AG}$. Hence, $\mathbf{Gb}$ is the least-squares solution minimizing the norm $\ \mathbf{Ax} - \mathbf{b}\ _2^2$.
$\{1, 4\}$	Minimum Norm Inverse	$\mathbf{G} \in \mathbf{A}^{\{1,4\}}$ iff $\mathbf{AGA} = \mathbf{A}$ and $(\mathbf{GA})^\top = \mathbf{GA}$. Hence, $\mathbf{Gb}$ is the minimum-norm solution of $\mathbf{Ax} = \mathbf{b}$ for all $\mathbf{b} \in \mathcal{R}(\mathbf{A})$.
$\{1, 2, 3, 4\}$	Moore-Penrose Inverse	The unique $\mathbf{G} \equiv \mathbf{A}^+ \in \mathbf{A}^{\{1,2,3,4\}}$ is the least-squares minimum-norm solution.

Source: Adapted from Campbell and Meyer (2009, Table 6.1, p. 96) [52] (pp. 91–119) as a summary of Rao and Mitra [33], pp. 44–71, Rao and Mitra [42], Ben-Israel and Greville [34], pp. 40–51, and Wang et al. [53], pp. 10–19, 33–42.

Starting from the early 1960s, Cline [47], Rao and Mitra [33], pp. 64–71, Meyer [54], Ben-Israel and Greville [34], pp. 175–200, Hartwig [55], Campbell and Meyer [52], pp. 53–61, Rao and Yanai [56], Tian [57], Rakha [58], Wang et al. [53], pp. 193–210, and Baksalary and Trenkler [59], among others, extended the formulas for $\mathbf{A}^+$ to partitioned (block) matrices, including the general column-wise, $\mathbf{A}_1 \in \mathbb{R}^{m \times n_1}$, $\mathbf{A}_2 \in \mathbb{R}^{m \times n_2}$, and $n_1 + n_2 = n$ — and, equivalently, row-wise, $\mathbf{A}_1 \in \mathbb{R}^{m_1 \times n}$, $\mathbf{A}_2 \in \mathbb{R}^{m_2 \times n}$, $m_1 + m_2 = m$, — cases (employed in the numerical stability analysis of the CLSP

estimator for $\mathbf{C}_{\text{canon}} = [\mathbf{C} \ \mathbf{S}]$ in the decomposition $\mathbf{A} = \mathbf{B}\mathbf{C}_{\text{canon}} + \mathbf{R}$ given $\mathbf{A} \in \mathbb{R}^{m \times n}$, $\mathbf{C}_{\text{canon}} \in \mathbb{R}^{k \times n}$, and $\mathbf{B} = \mathbf{A}\mathbf{C}_{\text{canon}}^{\dagger} = \mathbf{A}[\mathbf{C} \ \mathbf{S}]^{\dagger} \in \mathbb{R}^{m \times k}$:

$$[\mathbf{A}_1 \ \mathbf{A}_2]^{\dagger} = \left([\mathbf{A}_1 \ \mathbf{A}_2]^{\top} [\mathbf{A}_1 \ \mathbf{A}_2] \right)^{\dagger} [\mathbf{A}_1 \ \mathbf{A}_2]^{\top} = \begin{bmatrix} \mathbf{A}_1^{\top} \mathbf{A}_1 & \mathbf{A}_1^{\top} \mathbf{A}_2 \\ \mathbf{A}_2^{\top} \mathbf{A}_1 & \mathbf{A}_2^{\top} \mathbf{A}_2 \end{bmatrix}^{\dagger} \begin{bmatrix} \mathbf{A}_1^{\top} \\ \mathbf{A}_2^{\top} \end{bmatrix} \quad (2)$$

where $\text{rank}([\mathbf{A}_1 \ \mathbf{A}_2]) \leq \text{rank}(\mathbf{A}_1) + \text{rank}(\mathbf{A}_2)$, with equality iff $\mathcal{R}(\mathbf{A}_1) \cap \mathcal{R}(\mathbf{A}_2) = \{\mathbf{0}\}$, and strict inequality iff $\mathcal{R}(\mathbf{A}_1) \cap \mathcal{R}(\mathbf{A}_2) \neq \{\mathbf{0}\}$. For the original definitions and formulations, consult Rao and Mitra [33], pp. 64–66 and Ben-Israel and Greville [34], pp. 175–200.

To sum up, $\mathbf{G}$ (especially $\mathbf{A}^{\dagger}$), SVD, and regularization techniques — namely, Lasso (based on the ℓ_1 -norm), Tikhonov regularization (hereafter, referred to as Ridge regression, based on the ℓ_2 -norm), and Elastic Net (a convex combination of ℓ_1 and ℓ_2 norms) [4], pp. 477–479 — have become an integral part of estimation (model-fitting), where $\mathbf{A}^{\dagger}$ can be defined as a *minimum-norm best linear unbiased estimator* (MNBLUE), reducing to the classical BLUE in (over)determined cases. These methods have been applied in (a) natural sciences since the 1940s [32] (pp. 68–156) [36,51] [60], pp. 85–172 [61–63], and (b) econometrics since the 1960s [39,40] [64], pp. 36–157, 174–232 [65], pp. 61–106 [66], pp. 37–338 in both unconstrained and equality-constrained forms, thereby relaxing previous rank restrictions.

The incorporation of *inequality constraints* into optimization problems during the 1930s–1950s (driven by the foundational contributions of Kantorovich [67], Koopmans [7], and Dantzig [6], along with the later authors [5]) marked the next milestone in the development of convex optimization under the emerging framework of programming, namely for linear cases (LP) — formally, $\mathbf{x}^* = \arg \max_{\mathbf{x} \in \mathbb{R}^n} \mathbf{c}^{\top} \mathbf{x}$ subject to $\mathbf{A}\mathbf{x} \leq \mathbf{b}$, $\mathbf{x} \geq \mathbf{0}$ (with the dual $\lambda^* = \arg \min_{\lambda \in \mathbb{R}^m} \mathbf{b}^{\top} \lambda$ subject to $\mathbf{A}^{\top} \lambda \geq \mathbf{c}$), where $\mathbf{A} \in \mathbb{R}^{m \times n}$, $\mathbf{b} \in \mathbb{R}^m$, and $\mathbf{c} \in \mathbb{R}^n$ — and convex quadratic ones (QP) — formally, $\mathbf{x}^* = \arg \min_{\mathbf{x} \in \mathbb{R}^n} \left\{ \frac{1}{2} \mathbf{x}^{\top} \mathbf{Q} \mathbf{x} + \mathbf{c}^{\top} \mathbf{x} \right\}$ subject to $\mathbf{A}\mathbf{x} \leq \mathbf{b}$ (with the dual $\lambda^* = \arg \max_{\lambda \in \mathbb{R}_+^m} \left\{ -\frac{1}{2} (\mathbf{A}^{\top} \lambda + \mathbf{c})^{\top} \mathbf{Q}^{-1} (\mathbf{A}^{\top} \lambda + \mathbf{c}) - \mathbf{b}^{\top} \lambda \right\}$), where $\mathbf{Q} \in \mathbb{R}^{n \times n}$ is symmetric positive definite, $\mathbf{c} \in \mathbb{R}^n$, $\mathbf{A} \in \mathbb{R}^{m \times n}$, and $\mathbf{b} \in \mathbb{R}^m$ — [2] (pp. 146–213) (for a classification of programming, consult Figure 1) [6] (pp. 1–31) [1], pp. 355–597, with Dantzig's simplex and Karmarkar's interior-point algorithms being efficient solvers [68,69].

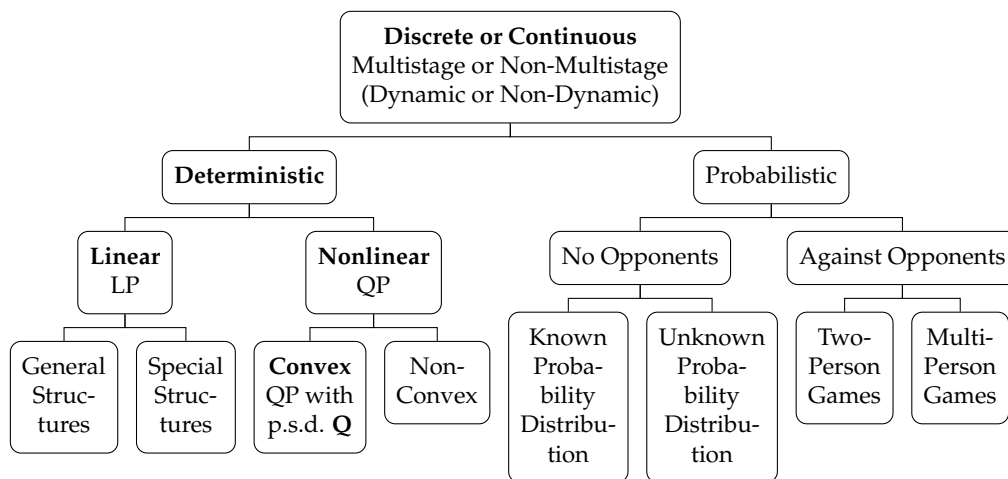


Figure 1. LP and QP in programming. Adapted from Dantzig (1990, Figure 1.4.1, p. 8) [6] (pp. 1–11).

In the late 1950s–1970s, the application of LP and QP extended to LS problems (which can be referred to as a *major generalization of mainstream convex optimization methods*; consult Boyd and Vandenberghe [2], pp. 291–349 for a comprehensive modern overview) with the pioneering works of Wagner [70] and Sielken and Hartley [71], expanded by Kiountouzis [72] and Sposito [73], who first substituted LS with LP-based (ℓ_1 -norm) least absolute (LAD) and (ℓ_∞ -norm) least maximum (LDP) deviations and derived unbiased (ℓ_p -norm, $1 < p < \infty$) estimators with non-unique solutions; whilst, in the early 1990s, Stone and Tovey [69] demonstrated algorithmic equivalence between LP

algorithms and iteratively reweighted LS. Judge and Takayama [74] and Mantel [75] reformulated (multiple) LS with inequality constraints as QP, and Lawson and Hanson [64], pp. 144–147 introduced, among others, the non-negative least squares (NNLS) method. These and further developments are reflected in the second step of the CLSP framework, where $\hat{\mathbf{z}} = \mathbf{G}\mathbf{b}$ is corrected under a regularization scheme inspired by Lasso ($\alpha = 0$), Ridge ($\alpha = 1$), and Elastic Net ($0 < \alpha < 1$) regularization, where $\hat{\mathbf{z}}^* = \arg \min_{\mathbf{z} \in \mathbb{R}^n} \left\{ (1 - \alpha) \|\mathbf{z} - \hat{\mathbf{z}}\|_1 + \alpha \|\mathbf{z} - \hat{\mathbf{z}}\|_2^2 \right\}$, subject to $\mathbf{A}\mathbf{z} = \mathbf{b}$.

The final frontier that motivates the formulation of a new framework is the class of *ill-posed problems* in Hadamard's sense [4] (pp. 143, 241), i.e., systems with no (i) feasibility, $\mathbf{x} \notin \mathcal{F} := \{\mathbf{x} \mid \forall_{i,j} g_i(\mathbf{x}) \leq 0, h_j(\mathbf{x}) = 0\}$, (ii) uniqueness, $\exists \mathbf{x}_1 \neq \mathbf{x}_2 : \mathbf{A}\mathbf{x}_1 = \mathbf{A}\mathbf{x}_2$, (iii) stability, $\forall \varepsilon > 0, \exists \delta > 0 : \|\mathbf{b} - \mathbf{b}_\delta\| < \delta \Rightarrow \|\mathbf{A}^\dagger \mathbf{b} - \mathbf{A}^\dagger \mathbf{b}_\delta\| \geq \varepsilon$, or formulation — as a well-posed LS, LP, or QP problem — of solution under mainstream assumptions, where $\mathbf{A} \in \mathbb{R}^{m \times n}$, $\mathbf{x} \in \mathbb{R}^n$, and $\mathbf{b} \in \mathbb{R}^m$, such as the ones described in Bolotov [15,16]: namely, (a) underdetermined, $\text{rank}(\mathbf{A}) < \min(m, n)$, and/or ill-conditioned linear systems, $\kappa(\mathbf{A}) = \|\mathbf{A}\|_2 \|\mathbf{A}^\dagger\|_2 \gg 1$, $\mathbf{A} \in \mathbb{R}^{m \times n}$, in cases of LS and LP, where solutions are either non-unique or highly sensitive to perturbations [76,77]; and (b) degenerate, $\#\{\text{active constraints at } \mathbf{x}\} > \dim(\mathbf{x}) = n$, with cycling behavior, $c^\top \mathbf{x}^{(k+1)} = c^\top \mathbf{x}^{(k)}$, and $\mathbf{x}^{(k+1)} \neq \mathbf{x}^{(k)}$, $\mathbf{x} \in \mathbb{R}^n$, in all LP and QP problems, where efficient methods, such as the simplex algorithm, may fail to converge [78,79]. Such problems have been addressed since the 1940s–1970s by (a) regularization (i.e., a method of perturbations) and (b) problem-specific algorithms [2] (pp. 455–630). For LS cases, the list includes, among others, Lasso, Ridge regularization, and Elastic Net [4] (pp. 477–479), constrained LS methods [80], restricted LS [81], LP-based sparse recovery approaches [77], Gröbner bases [82], as well as derivatives of eigenvectors and Nelson-type, BD, QR- and SVD-based algorithms [83]. For LS and QP — among others, Wolfe's regularization-based technique for simplex [84] and its modifications [78,79], Dax's LS-based steepest-descent algorithm [85], a primal-dual NNLS-based algorithm (LSPD) [86], as well as modifications of the ellipsoid algorithm [87]. A major early attempt to unify (NN)LS, LP, and QP within a single framework to solve both well- and ill-posed convex optimization problems was undertaken by Übi [88,89,90,91,92], who proposed a series of problem-specific algorithms (VD, INE, QPLS, and LS1). Figure 2 compares all the mentioned methods by functionality.

Problem Class/ Method		LS, based on $\mathbf{A}^{-1}$	LS, based on $\mathbf{A}^\dagger$ or $\mathbf{A}_S^\dagger$	LP/QP	Regular- ization	Problem- specific
Well-posed unconstrained or equality- constrained	Full-rank	•	•	•	•	•
	Overdetermined	•	•	•	•	•
	Underdetermined		•	•	•	•
Well-posed with inequality constraints	Full-rank			•	•	•
	Overdetermined			•	•	•
	Underdetermined			•	•	•
Ill-posed unconstrained or constrained	Full-rank				•	•
	Overdetermined				•	•
	Underdetermined				•	•

Figure 2. Applicability of selected types of convex optimization methods across problem classes.

Therefore, to accommodate all considered problem classes, a *modular two-step CLSP* consists of (first, compulsory, step) minimum-norm LS estimation, based on $\mathbf{A}^\dagger$ or $\mathbf{A}_S^\dagger$, followed by a (second, optional, step) **regularization-based correction**, contributing to the generalization of convex optimization methods by eliminating the need for problem-specific algorithms.

3. Construction and Formalization of the CLSP Estimator

Based on the citation and reference counts, as well as in-group relevance (i.e., the number of citations within the examined collection) of the selected sources (consult Figure 3), this text incorporates the seminal studies on the topic, with the exception of historical seminar proceedings (e.g., Nashed) and a few secondary contributions by Ben-Israel, Sergienko, and Grafarend (Lawson and Hanson,

erroneously indicated as missing due to incomplete monograph indexing, are already included in the bibliography as a 1995 reprint of the 1976 edition [64]). Hence, it can be claimed that the CLSP framework is grounded in the general algorithmic logic of Wolfe [84] and Dax [85], as further formalized by Osborne [78], who were among the first to introduce an additional estimation step into simplex-based solvers.

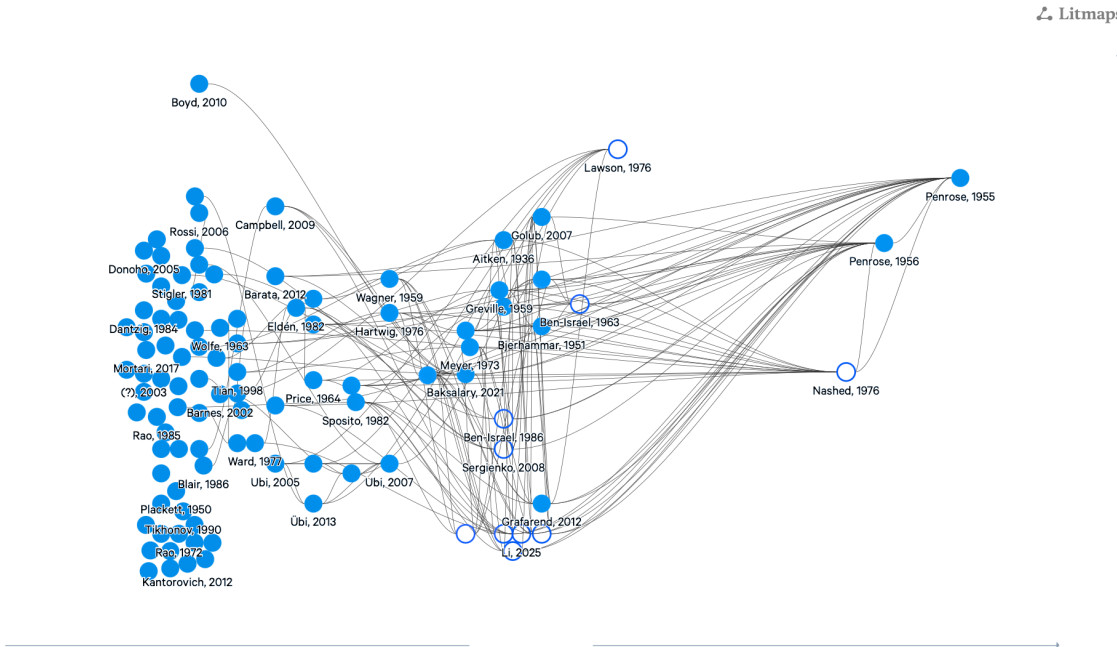


Figure 3. Map of citation and reference counts for the above-cited sources, where hollow circles are suggested items. Monographs not indexed in the database are denoted as (?). Adapted from Litmaps (Shared Citations & References), available at <https://www.litmaps.com/> (accessed on June 30, 2025).

The proposed CLSP estimator is *modular* — i.e., it comprises selectively omissible and repeatable actions for enhanced flexibility — and consists of *two steps*: (a) the first, grounded in the theory of (constrained) pseudoinverses as summarized in Rao and Mitra [33], Ben-Israel and Greville [34], Lawson and Hanson [64], and Wang et al. [53], is iterable to refine the solution and serves as a mandatory baseline in cases where the second step becomes infeasible (due to mutually inconsistent constraints) — as exemplified by Whiteside et al. [93] for regression-derived constraints and by Blair [76,94] in systems with either too few or too many effective constraints; — and (b) the second, grounded in the theory of regularization and convex programming as summarized in Tikhonov et al. [32], Gentle [4], Nocedal and Wright [1], and Boyd and Vandenberghe [2], provides an optional correction of the first step, drawing on the logic of Lasso, Ridge, and Elastic Net to address ill-posedness and compensate for constraint violations or residual approximation errors resulting from the first-step minimum-norm LS estimate. This structure ensures that CLSP *is able to yield a unique solution* under a strictly convex second-step correction and *extends the classical LS framework in both scope and precision*. The modular algorithmic flow is illustrated in Figure 4.

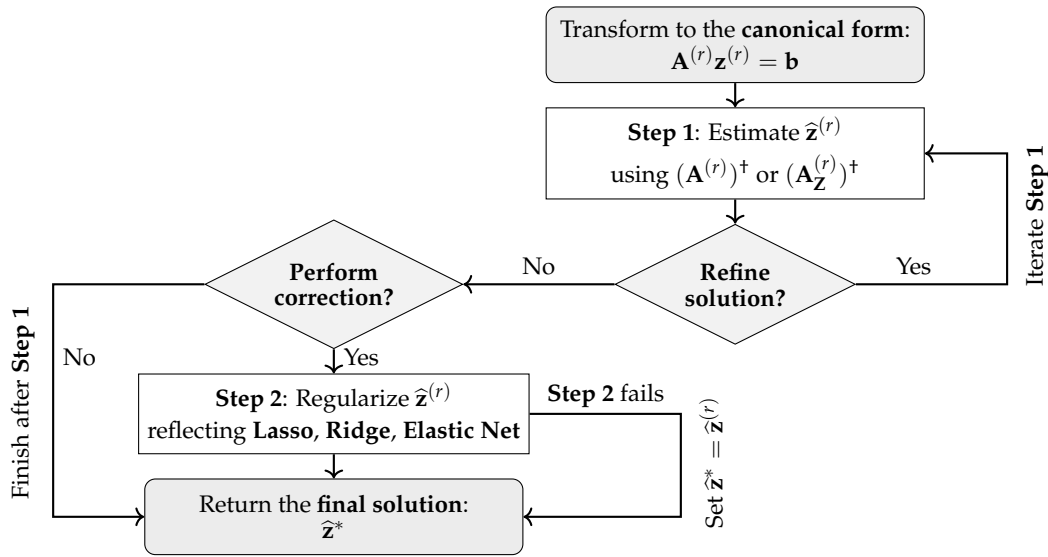


Figure 4. Algorithmic flow of the CLSP Estimator. Decision nodes (diamonds) indicate modularity.

The first process block in the CLSP algorithmic flow denotes the transformation of the initial problem — $\hat{\mathbf{x}}_M^* = \text{Proj}_{\mathbf{x}_M} \left\{ \mathbf{Q}_X [\mathbf{x}_L^M] \in \mathbb{R}^p : \forall_{1 \leq i \leq s} g_i(\mathbf{x}_M, \mathbf{x}_L) \geq \gamma_i, \forall_j h_j(\mathbf{x}_M, \mathbf{x}_L) = \eta_j \right\}$, where $\text{Proj}_{\mathbf{x}_M}$ is the projection onto the $\mathbf{x}_M$ coordinates, $\mathbf{x}_M \in \mathbb{R}^{p-l}$ is the vector of model (target) variables to be estimated ($\hat{\mathbf{x}}_M^*$), $\mathbf{x}_L \in \mathbb{R}^l$, $l \geq 0$, is a vector of latent variables that appear only in the constraint functions $g_i(\cdot)$ and $h_j(\cdot)$, such that $\mathbf{Q}_X [\mathbf{x}_L^M] = \mathbf{x} \in \mathbb{R}^p$, $\mathbf{Q}_X$ is the permutation matrix, and $g_i(\cdot) \geq \gamma_i$ and $h_j(\cdot) = \eta_j$ are the linear(ized) inequality and equality constraints, $\forall_{i,j} \gamma_i, \eta_j \in \mathbf{b}$ — to the **canonical form** $\mathbf{A}^{(r)} \mathbf{z}^{(r)} = \mathbf{b}$ (a term first introduced in LP [6] (pp. 75–81)), where (a) $\mathbf{A}^{(r)} \in \mathbb{R}^{m \times n}$ is the block *design matrix* consisting of a $\forall_i g_i(\cdot)$ -and- $\forall_j h_j(\cdot)$ *constraints matrix* $\mathbf{C} \in \mathbb{R}^{k \times p}$, a *model matrix* $\mathbf{M} = \mathbf{Q}_M^L \begin{bmatrix} \mathbf{M}_{\text{partial}} \\ \mathbf{0} \end{bmatrix} \mathbf{Q}_M^R \in \mathbb{R}^{(m-k) \times p}$, a *sign slack matrix* $\mathbf{S} = \mathbf{Q}_S^L \begin{bmatrix} \mathbf{S}_s \subseteq \forall_{\sigma_1, \dots, \sigma_s \in \{-1, +1\}} \text{diag}(\sigma_1, \dots, \sigma_s) \\ \mathbf{0} \end{bmatrix} \mathbf{Q}_S^R \in \mathbb{R}^{k \times (n-p)}$, and either a *zero matrix* $\mathbf{0}$ (if $r = 1$) or a *reverse-sign* — i.e., based on $-\text{sign}(\cdot)$ — *slack matrix* $\mathbf{S}_{\text{residual}} = \begin{bmatrix} \mathbf{0} & \text{diag}(-\text{sign}((\mathbf{b} - \mathbf{A}^{(r-1)} \mathbf{z}^{(r-1)})_{k+1:m})) \end{bmatrix} \in \mathbb{R}^{(m-k) \times (n-p)}$ (provided $r > 1$):

$$\mathbf{A}^{(r)} = \begin{bmatrix} \mathbf{C} & \mathbf{S} \\ \mathbf{M} & \begin{cases} \mathbf{0}, & \text{initial iteration, } r = 1 \\ \mathbf{S}_{\text{residual}}, & \text{subsequent iterations, } r > 1 \end{cases} \end{bmatrix} \quad (3)$$

and (b) $\mathbf{z}^{(r)} = \begin{bmatrix} \mathbf{x}^{(r)} \\ (\mathbf{y}^{(r)})^* \end{bmatrix} \in \mathbb{R}^n$ is the *full vector of unknowns*, comprising (model) $\mathbf{x}^{(r)} \in \mathbb{R}^p$ and (slack) $(\mathbf{y}^{(r)})^* = \begin{bmatrix} \mathbf{y}^{(r)} \\ \mathbf{y}_{\text{residual}}^{(r)} \end{bmatrix} \in \mathbb{R}^{n-p}$, $y_i^{(r)} \geq 0$ and $y_j^{(r)} = 0$ — an equivalent problem [2] (pp. 130–132), where $g_i(\cdot)^a = g_i(\cdot) + \mathbf{S}_{ii} y_i^{(r)} = 0$ and $h_j(\cdot)^a = h_j(\cdot)$. The constraint functions $g_i(\cdot)^a$ and $h_j(\cdot)^a$ must be linear(ized), so that $\mathbf{C}, \mathbf{S} \subseteq \mathbf{A}^{(r)} \Leftrightarrow \mathbf{C} \mathbf{x}^{(r)} + \mathbf{S} (\mathbf{y}^{(r)})^* = \mathbf{b}_{1:k}$, $k \leq m$.

In practical implementations, for problems with $(m > m_E) \vee (n > n_E)$, where E is an estimability limit, the partitioning $\mathbf{A}^{(r)} = \left\{ \mathbf{A}_{\text{subset}, i_E, j_E}^{(r)} \right\}$, $m_{\text{subset}} = m_{\text{reduced}} - 1$, $n_{\text{subset}} = n_{\text{reduced}} - 1$, $i_E \in \{1, \dots, \lceil \frac{m}{m_{\text{subset}}} \rceil\}$, and $j_E \in \{1, \dots, \lceil \frac{n}{n_{\text{subset}}} \rceil\}$, can be used, leading to a total of $\left(\lceil \frac{m}{m_{\text{subset}}} \rceil \cdot \lceil \frac{n}{n_{\text{subset}}} \rceil \right)$ $\mathbf{A}_{\text{reduced}}^{(r)} \in \mathbb{R}^{m_{\text{reduced}} \times n_{\text{reduced}}}$, composed of $\mathbf{A}_{\text{subset}}^{(r)}$ and $\sum_{i \notin \mathcal{I}_{\text{subset}}} a_{i,j}^{(r)}$, $i \notin \mathcal{I}_{\text{subset}} \subset \{1, \dots, m\}$, $j \notin \mathcal{J}_{\text{subset}} \subset \{1, \dots, n\}$, with a total of $\left(\lceil \frac{m}{m_{\text{subset}}} \rceil \right) \mathbf{b}_r \in \mathbb{R}^{m_{\text{reduced}}}$:

$$\mathbf{A}_{\text{reduced}}^{(r)} = \mathbf{Q}_A^T \begin{bmatrix} \mathbf{A}_{\mathcal{I}_{\text{subset}}, \mathcal{J}_{\text{subset}}}^{(r)} & \sum_{j \notin \mathcal{J}_{\text{subset}}} a_{i,j}^{(r)} \\ \sum_{i \notin \mathcal{I}_{\text{subset}}} a_{i,j}^{(r)} & \sum_{i \notin \mathcal{I}_{\text{subset}}} \sum_{j \notin \mathcal{J}_{\text{subset}}} a_{i,j}^{(r)} \end{bmatrix} \mathbf{Q}_A^T, \quad \mathbf{b}_r = \mathbf{Q}_b^T \begin{bmatrix} \mathbf{b}_{\mathcal{I}_{\text{subset}}} \\ \sum_{i \notin \mathcal{I}_{\text{subset}}} b_i \end{bmatrix} \quad (4)$$

$\sum a_{i,j}^{(r)}$ are then treated as slack rows/columns, i.e., incorporated into $\mathbf{S}$, to act as balancing adjustments since the reduced model, by construction, tends to inflate the estimates $\hat{a}_{i_E, j_E}$.

The second process block, i.e., the **first step of the CLSP estimator**, denotes obtaining an (iterated if $r > 1$) estimate $\hat{\mathbf{z}}^{(r)} = \mathbf{G}^{(r)} \mathbf{b}$ (alternatively, $\left(\left\lceil \frac{m}{m_{\text{subset}}} \right\rceil \cdot \left\lceil \frac{n}{n_{\text{subset}}} \right\rceil\right)$ -times $\hat{\mathbf{z}}_{\text{reduced}}^{(r)} = \mathbf{G}_{\text{reduced}}^{(r)} \mathbf{b}_r$ with $\hat{\mathbf{z}}_{i_E, j_E}^{(r)} = \left\{ \hat{\mathbf{z}}_{\text{reduced}, 1:m_{\text{subset}}, 1:n_{\text{subset}}}^{(r), (i_E, j_E)} \right\}$) through the (constrained) Bott-Duffin inverse $\mathbf{G}^{(r)} \equiv (\mathbf{A}_{\mathbf{Z}}^{(r)})^{\dagger} \in \mathbb{R}^{n \times m}$ [51] [53] (pp. 49–64), defined on a subspace specified by a symmetric idempotent matrix $\mathbf{Z} \in \mathbb{R}^{n \times n}$ (regulating solution-space exclusion through binary entries 0/1, such that $\mathbf{P}_{\mathbf{Z}}^L = \mathbf{Z}$ restricts the domain of estimated variables to the selected subspace, while $\mathbf{P}_{\mathbf{Z}}^R = \mathbf{I}_m$ leaves the data vector, i.e., the input $\mathbf{b}$, unprojected). $\mathbf{G}^{(r)} \equiv (\mathbf{A}_{\mathbf{Z}}^{(r)})^{\dagger}$ reduces to the Moore-Penrose pseudoinverse $\mathbf{G}^{(r)} \equiv (\mathbf{A}^{(r)})^{\dagger} \in \mathbb{R}^{n \times m}$ when $\mathbf{Z} = \mathbf{I}_n$, and equals the standard inverse $\mathbf{G}^{(r)} \equiv (\mathbf{A}^{(r)})^{-1}$ iff $\mathbf{Z} = \mathbf{I}_n \wedge \text{rank}(\mathbf{A}^{(r)}) = m = n$:

$$\hat{\mathbf{z}}^{(r)} = (\mathbf{A}_{\mathbf{Z}}^{(r)})^{\dagger} \mathbf{b} = \left[\left(\mathbf{Z}(\mathbf{A}^{(r)})^{\top} \mathbf{A}^{(r)} \mathbf{Z} \right)^{\dagger} \mathbf{Z}(\mathbf{A}^{(r)})^{\top} \right] \mathbf{b}, \quad \mathbf{Z}^2 = \mathbf{Z} = \mathbf{Z}^{\top} \quad (5)$$

The number of iterations (r) increases until a termination condition is met, such as error $\left| \frac{1}{\sqrt{m}} \left\| (\mathbf{b} - \mathbf{A}^{(r)} \hat{\mathbf{z}}^{(r)}) \right\|_2 - \frac{1}{\sqrt{m}} \left\| (\mathbf{b} - \mathbf{A}^{(r-1)} \hat{\mathbf{z}}^{(r-1)}) \right\|_2 \right| < \varepsilon$ or $r > r_{\text{limit}}$, whichever occurs first. In practical implementations, $\hat{\mathbf{z}}^{(r)}$ is efficiently computed using the SVD (see **Appendix A.1**).

The third process block, i.e., the **second step of the CLSP estimator**, denotes obtaining a regularized solution $\hat{\mathbf{z}}^*$ (alternatively, $\left(\left\lceil \frac{m}{m_{\text{subset}}} \right\rceil \cdot \left\lceil \frac{n}{n_{\text{subset}}} \right\rceil\right)$ -times $\hat{\mathbf{z}}_{\text{reduced}}^*$ with $\hat{\mathbf{z}}_{i_E, j_E}^* = \left\{ \hat{\mathbf{z}}_{\text{reduced}, 1:m_{\text{subset}}, 1:n_{\text{subset}}}^{*, (i_E, j_E)} \right\}$) — unless the algorithmic flow terminates after Step 1, in which case $\hat{\mathbf{z}}^* = \hat{\mathbf{z}}^{(r)}$ — by penalizing deviations from the (iterated if $r > 1$) estimate $\hat{\mathbf{z}}^{(r)}$ reflecting Lasso ($\alpha = 0 \Leftrightarrow \hat{\mathbf{z}}^* = \hat{\mathbf{z}}^{(r)} + \hat{\mathbf{u}}$, where $\hat{\mathbf{u}} = \arg \min_{\mathbf{u}} \|\mathbf{u}\|_1$, s.t. $\mathbf{A}^{(r)}(\hat{\mathbf{z}}^{(r)} + \mathbf{u}) = \mathbf{b}$), Ridge ($\alpha = 1 \Leftrightarrow \hat{\mathbf{z}}^* = \hat{\mathbf{z}}^{(r)} + (\mathbf{A}^{(r)})^{\dagger} \mathbf{A}^{(r)} \hat{\mathbf{u}}$, where $(\mathbf{A}^{(r)})^{\dagger} \mathbf{A}^{(r)} = \mathbf{P}_{\mathcal{R}(\mathbf{A}^{\top})}$, a projector equal to $\mathbf{I}$ iff $\mathbf{A}^{(r)}$ is of full column rank, and $\hat{\mathbf{u}} = \arg \min_{\mathbf{u}} \|\mathbf{u}\|_2^2$, s.t. $\mathbf{A}^{(r)}(\hat{\mathbf{z}}^{(r)} + \mathbf{u}) = \mathbf{b}$), and Elastic Net ($0 < \alpha < 1 \Leftrightarrow \hat{\mathbf{z}}^* = \hat{\mathbf{z}}^{(r)} + (1 - \alpha)\hat{\mathbf{u}} + \alpha(\mathbf{A}^{(r)})^{\dagger} \mathbf{A}^{(r)} \hat{\mathbf{u}}$, where $\hat{\mathbf{u}} = \arg \min_{\mathbf{u}} \{(1 - \alpha)\|\mathbf{u}\|_1 + \alpha\|\mathbf{u}\|_2^2\}$, s.t. $\mathbf{A}^{(r)}(\hat{\mathbf{z}}^{(r)} + \mathbf{u}) = \mathbf{b}$) [4] (pp. 477–479), [2] (pp. 306–310):

$$\hat{\mathbf{z}}^* = \arg \min_{\mathbf{z} \in \mathbb{R}^n} \left\{ (1 - \alpha) \left\| \mathbf{z} - \hat{\mathbf{z}}^{(r)} \right\|_1 + \alpha \left\| \mathbf{z} - \hat{\mathbf{z}}^{(r)} \right\|_2^2 \right\}, \quad \alpha \in [0, 1], \quad \text{s.t. } \mathbf{A}^{(r)} \mathbf{z} = \mathbf{b} \quad (6)$$

where the parameter α may be selected arbitrarily or determined from the input $\mathbf{b}$ using an appropriate criterion, such as minimizing prediction error via cross-validation, satisfying model-selection heuristics (e.g., AIC or BIC), or optimizing residual metrics under known structural constraints. Alternatively, α may be fixed at benchmark values, e.g., $\alpha \in \{0, \frac{1}{2}, 1\}$, or based on an error rule, e.g., $\alpha = \min \left(\frac{\left(\frac{1}{\sqrt{m}} \left\| (\mathbf{b} - \mathbf{A}^{(r)} \hat{\mathbf{z}}^*) \right\|_2 \right)_{\alpha=0}}{\left(\frac{1}{\sqrt{m}} \left\| (\mathbf{b} - \mathbf{A}^{(r)} \hat{\mathbf{z}}^*) \right\|_2 \right)_{\alpha=0} + \left(\frac{1}{\sqrt{m}} \left\| (\mathbf{b} - \mathbf{A}^{(r)} \hat{\mathbf{z}}^*) \right\|_2 \right)_{\alpha=1} + \varepsilon}, 1 \right)$, $\varepsilon \rightarrow 0^+$. The final value $\hat{\mathbf{z}}^*$ is obtained by solving a convex programming problem, which, in practical implementations, $\hat{\mathbf{z}}^*$ is efficiently computed with the help of numerical solvers.

This formalization of the process blocks ensures that the CLSP estimator yields a unique final solution $\hat{\mathbf{z}}^*$ under regularization weight $\alpha \in (0, 1]$ owing to the strict convexity of the objective function (consult Theorem 1 for a proof). For $\alpha = 0$, the solution remains convex but may not be unique unless additional structural assumptions are imposed. In the specific case of $\alpha = 1$, the estimator reduces to the Minimum-Norm Best Linear Unbiased Estimator (MNBLUE), or, alternatively, to a block-wise $m_{\text{subset}} \times n_{\text{subset}}$ -MNBLUE across $\left(\left\lceil \frac{m}{m_{\text{subset}}} \right\rceil \cdot \left\lceil \frac{n}{n_{\text{subset}}} \right\rceil\right)$ reduced problems (for a proof, consult Theorem 2 and its corollaries).

Theorem 1. Let $\hat{\mathbf{z}}^{(r)} = (\mathbf{A}_{\mathbf{Z}}^{(r)})^{\dagger} \mathbf{b}$ be the r -iterated estimate obtained in the first step of the CLSP algorithm, and let $\hat{\mathbf{z}}^*$ be the second-step correction obtained through convex programming by solving $\hat{\mathbf{z}}^* = \arg \min_{\mathbf{z} \in \mathbb{R}^n} \left\{ (1 - \alpha) \left\| \mathbf{z} - \hat{\mathbf{z}}^{(r)} \right\|_1 + \alpha \left\| \mathbf{z} - \hat{\mathbf{z}}^{(r)} \right\|_2^2 \right\}$, $\alpha \in (0, 1]$, s.t. $\mathbf{A}^{(r)} \mathbf{z} = \mathbf{b}$ (i.e., the regularization parameter excludes a pure Lasso-type correction), then the final solution $\hat{\mathbf{z}}^*$ is unique.

Proof. Let $\hat{\mathbf{z}}^{(r)} = (\mathbf{A}_Z^{(r)})^\dagger \mathbf{b}$ denote the linear estimator obtained via the Bott-Duffin inverse, defined on a subspace determined by a symmetric idempotent matrix $\mathbf{Z}$, and producing a conditionally unbiased estimate over that subspace. The Bott-Duffin inverse $(\mathbf{A}_Z^{(r)})^\dagger$ is given by $(\mathbf{A}_Z^{(r)})^\dagger = (\mathbf{Z}(\mathbf{A}^{(r)})^\top \mathbf{A}^{(r)} \mathbf{Z})^\dagger \mathbf{Z}(\mathbf{A}^{(r)})^\top$, and the estimate $\hat{\mathbf{z}}^{(r)}$ is unique if $\text{rank}(\mathbf{A}^{(r)} \mathbf{Z}) = \dim(\mathcal{R}(\mathbf{Z}))$. In this case, $\hat{\mathbf{z}}^{(r)}$ is the unique minimum-norm solution in $\mathcal{R}(\mathbf{Z})$. In all other cases, the solution set is affine and given by $\hat{\mathbf{z}}^{(r)} = (\mathbf{A}_Z^{(r)})^\dagger \mathbf{b} + [\ker(\mathbf{A}^{(r)} \mathbf{Z}) \cap \mathcal{R}(\mathbf{Z})]$, where the null-space component represents degrees of freedom not fixed by the constraint $\mathbf{A}^{(r)} \mathbf{Z}$, hence $\hat{\mathbf{z}}^{(r)}$ is not unique. At the same time, the second-step estimate $\hat{\mathbf{z}}^*$ is obtained by minimizing the function $(1 - \alpha) \|\mathbf{z} - \hat{\mathbf{z}}^{(r)}\|_1 + \alpha \|\mathbf{z} - \hat{\mathbf{z}}^{(r)}\|_2^2$ over the affine (hence convex and closed) constraint set $\mathcal{F}^{(r)} = \{\mathbf{z} \in \mathbb{R}^n \mid \mathbf{A}^{(r)} \mathbf{z} = \mathbf{b}\}$. Under $\alpha \in (0, 1]$, the quadratic term $\alpha \|\mathbf{z} - \hat{\mathbf{z}}^{(r)}\|_2^2$ contributes a positive-definite Hessian $2\alpha \mathbf{I} \succ 0$ to the objective function, making it a strictly convex parabola-like curvature, and, given $\mathcal{F}^{(r)} \neq \emptyset$, the minimizer $\mathbf{z}^*$ exists and is unique. Therefore, the CLSP estimator with $\alpha \in (0, 1]$ yields a unique $\hat{\mathbf{z}}^*$. $\square$

Corollary 1. Let the final solution be $\hat{\mathbf{z}}^* = \hat{\mathbf{z}}^{(r)} = (\mathbf{A}_Z^{(r)})^\dagger \mathbf{b}$, where $(\mathbf{A}_Z^{(r)})^\dagger$ denotes the Bott-Duffin inverse on a subspace defined by a symmetric idempotent matrix $\mathbf{Z}$. Then (one-step) $\hat{\mathbf{z}}^*$ is unique iff $\text{rank}(\mathbf{A}^{(r)} \mathbf{Z}) = \dim(\mathcal{R}(\mathbf{Z}))$. Else, the solution set is affine and given by $\hat{\mathbf{z}}^* \in (\mathbf{A}_Z^{(r)})^\dagger \mathbf{b} + [\ker(\mathbf{A}^{(r)} \mathbf{Z}) \cap \mathcal{R}(\mathbf{Z})]$, which implies that the minimizer $\mathbf{z}^*$ is not unique, i.e., $\#\arg \min(\cdot) > 1$.

Corollary 2. Let $\hat{\mathbf{z}}^*$ be the final solution obtained in the second step of the CLSP estimator by solving $\hat{\mathbf{z}}^* = \arg \min_{\mathbf{z} \in \mathbb{R}^n} \left\{ \|\mathbf{z} - \hat{\mathbf{z}}^{(r)}\|_1 \right\}$, s.t. $\mathbf{A}^{(r)} \mathbf{z} = \mathbf{b}$ (i.e., the regularization parameter is set to $\alpha = 0$, corresponding to a Lasso correction). Then the solution $\hat{\mathbf{z}}^*$ exists and the problem is convex. The solution is unique iff the following two conditions are simultaneously satisfied: (1) the affine constraint set $\mathcal{F}^{(r)} = \{\mathbf{z} \in \mathbb{R}^n \mid \mathbf{A}^{(r)} \mathbf{z} = \mathbf{b}\}$ intersects the subdifferential of $\|\mathbf{z} - \hat{\mathbf{z}}^{(r)}\|_1$ at exactly one point, and (2) the objective function is strictly convex on $\mathcal{F}^{(r)}$, which holds when $\mathcal{F}^{(r)}$ intersects the interior of at most one orthant in $\mathbb{R}^n$. In all other cases, the minimizer $\mathbf{z}^*$ is not unique, and the set of optimal solutions forms a convex subset of the feasible region $\mathcal{F}^{(r)}$; that is, the final solution $\hat{\mathbf{z}}^* \in \arg \min_{\mathbf{z} \in \mathcal{F}^{(r)}} \|\mathbf{z} - \hat{\mathbf{z}}^{(r)}\|_1$, where $\arg \min(\cdot)$ is convex and $\#\arg \min(\cdot) > 1$.

Theorem 2. Let $\hat{\mathbf{z}}^{(r)} = (\mathbf{A}_Z^{(r)})^\dagger \mathbf{b}$ be the r -iterated estimate obtained in the first step of the CLSP algorithm, and let $\hat{\mathbf{z}}^*$ be the second-step correction obtained through convex programming by solving $\hat{\mathbf{z}}^* = \arg \min_{\mathbf{z} \in \mathbb{R}^n} \left\{ \|\mathbf{z} - \hat{\mathbf{z}}^{(r)}\|_2^2 \right\}$, s.t. $\mathbf{A}^{(r)} \mathbf{z} = \mathbf{b}$ (i.e., the regularization parameter is $\alpha = 1$, corresponding to a Ridge correction), then $\hat{\mathbf{z}}^*$ is the Minimum-Norm Best Linear Unbiased Estimator (MNBLUE) of $\mathbf{z}$ under the linear(ized) constraints set by the canonical form of the CLSP.

Proof. Let $\hat{\mathbf{z}}^{(r)} = (\mathbf{A}_Z^{(r)})^\dagger \mathbf{b}$ denote the linear estimator obtained via the Bott-Duffin inverse, defined on a subspace determined by a symmetric idempotent matrix $\mathbf{Z}$, and producing a conditionally unbiased estimate over that subspace. The Bott-Duffin inverse is $(\mathbf{A}_Z^{(r)})^\dagger = (\mathbf{Z}(\mathbf{A}^{(r)})^\top \mathbf{A}^{(r)} \mathbf{Z})^\dagger \mathbf{Z}(\mathbf{A}^{(r)})^\top$ and the estimate $\hat{\mathbf{z}}^{(r)}$ is unique if $\text{rank}(\mathbf{A}^{(r)} \mathbf{Z}) = \dim(\mathcal{R}(\mathbf{Z}))$. Given the linear model derived from the canonical form $\mathbf{b} = \mathbf{A}^{(r)} \mathbf{z}^{(r)} + \boldsymbol{\varepsilon}$, where $\boldsymbol{\varepsilon}$ are residuals with $\mathbb{E}[\boldsymbol{\varepsilon}] = \mathbf{0}$, it follows that $\mathbb{E}[\hat{\mathbf{z}}^{(r)}] = \mathbb{E}[(\mathbf{A}_Z^{(r)})^\dagger \mathbf{b}] = (\mathbf{A}_Z^{(r)})^\dagger \mathbb{E}[\mathbf{b}] = (\mathbf{A}_Z^{(r)})^\dagger \mathbf{A}^{(r)} \mathbf{z}^{(r)}$. Substituting $(\mathbf{A}_Z^{(r)})^\dagger$ yields $\mathbb{E}[\hat{\mathbf{z}}^{(r)}] = (\mathbf{Z}(\mathbf{A}^{(r)})^\top \mathbf{A}^{(r)} \mathbf{Z})^\dagger \mathbf{Z}(\mathbf{A}^{(r)})^\top \mathbf{A}^{(r)} \mathbf{z}^{(r)}$. By the generalized projection identity $(\mathbf{A}_Z^{(r)})^\dagger \mathbf{A}^{(r)} \mathbf{Z} = \mathbf{Z}$, one obtains $\mathbb{E}[\hat{\mathbf{z}}^{(r)}] = \mathbf{Z} \mathbf{z}^{(r)}$, which proves conditional unbiasedness on $\mathcal{R}(\mathbf{Z})$ (and full unbiasedness if $\mathbf{Z} = \mathbf{I}$). Subsequently, the second-step estimate $\hat{\mathbf{z}}^*$ is obtained by minimizing the squared Euclidean distance between $\hat{\mathbf{z}}^{(r)}$ and $\mathbf{z}$, subject to the affine constraint $\mathbf{A}^{(r)} \mathbf{z} = \mathbf{b}$, which corresponds, in geometric and algebraic terms, to an orthogonal projection of $\hat{\mathbf{z}}^{(r)}$ onto the affine (hence convex and closed) subspace $\mathcal{F}^{(r)} = \{\mathbf{z} \in \mathbb{R}^n \mid \mathbf{A}^{(r)} \mathbf{z} = \mathbf{b}\}$. By Theorem 1, the unique minimizer $\mathbf{z}^*$ exists and is given

explicitly by $\mathbf{z}^* = \hat{\mathbf{z}}^{(r)} + (\mathbf{A}^{(r)})^\dagger (\mathbf{b} - \mathbf{A}^{(r)}\hat{\mathbf{z}}^{(r)})$. This expression satisfies the first-order optimality condition of the convex program and ensures that $\mathbf{z}^*$ is both feasible and closest (in the examined ℓ_2 -norm) to $\hat{\mathbf{z}}^{(r)}$ among all solutions in $\mathcal{F}^{(r)}$. Therefore, $\mathbf{z}^*$ is (1) linear (being an affine transformation of a linear estimator), (2) unbiased (restricted to the affine feasible space), (3) minimum-norm (by construction), and (4) unique (by strict convexity). By the generalized Gauss-Markov theorem [33] (Chapter 7, Definition 4, p. 139), [34] (Chapter 8, Section 3.2, Theorem 2, p. 287), it follows that $\hat{\mathbf{z}}^*$ is the *Minimum-Norm Best Linear Unbiased Estimator* (MNBLUE) under the set of linear(ized) constraints $\mathbf{A}^{(r)}\mathbf{z} = \mathbf{b}$. $\square$

Corollary 3. Let the canonical system $\mathbf{A}^{(r)}\mathbf{z}^{(r)} = \mathbf{b}$ be consistent and of full column rank, $\mathbf{A}^{(r)} \in \mathbb{R}^{m \times n}$, $m \geq n$, and $\text{rank}(\mathbf{A}^{(r)}) = n$. Suppose that the CLSP algorithm terminates after Step 1 and that $\mathbf{Z} = \mathbf{I}_n$, such that $\hat{\mathbf{z}}^* = \hat{\mathbf{z}}^{(r)} = (\mathbf{A}^{(r)})^\dagger \mathbf{b} = ((\mathbf{A}^{(r)})^\top \mathbf{A}^{(r)})^{-1} (\mathbf{A}^{(r)})^\top \mathbf{b}$. Then, provided the classical linear model $\mathbf{b} = \mathbf{A}^{(r)}\mathbf{z}^{(r)} + \varepsilon$, where $\mathbb{E}[\varepsilon] = \mathbf{0}$ and $\text{Cov}(\varepsilon) = \sigma^2 \mathbf{I}_m$, holds, the CLSP estimator is equivalent to OLS while $\hat{\mathbf{z}}^*$ is the *Best Linear Unbiased Estimator* (BLUE).

Corollary 4. Let $\hat{\mathbf{z}}_{\text{reduced}}^{(r)} = \mathbf{A}_{\mathbf{z}_{\text{reduced}}, \text{reduced}}^\dagger \mathbf{b}_r$ be the r -iterated estimate obtained via the Bott-Duffin inverse, defined on a subspace determined by a symmetric idempotent matrix $\mathbf{Z}_{\text{reduced}}$, and producing a conditionally unbiased estimate over that subspace. Subsequently, the estimate $\hat{\mathbf{z}}_{\text{reduced}}^*$ is obtained by minimizing the squared Euclidean distance between $\hat{\mathbf{z}}_{\text{reduced}}^{(r)}$ and $\mathbf{z}_{\text{reduced}}$, subject to the affine constraint $\mathbf{A}_{\text{reduced}}^{(r)} \mathbf{z}_{\text{reduced}} = \mathbf{b}_r$, such that $\hat{\mathbf{z}}_{i_E, j_E}^* = \left\{ \hat{\mathbf{z}}_{\text{reduced}, 1:m_{\text{subset}}, 1:n_{\text{subset}}}^*, (i_E, j_E) \right\}$. Then each $\hat{\mathbf{z}}_{\text{reduced}}^*$ is the *Minimum-Norm Best Linear Unbiased Estimator* (MNBLUE) under the linear(ized) constraints defined by the canonical form of each reduced system of the CLSP, and $\hat{\mathbf{z}}^*$ is a block-wise $m_{\text{subset}} \times n_{\text{subset}}$ -MNBLUE corresponding to each of the $\left(\left\lceil \frac{m}{m_{\text{subset}}} \right\rceil \cdot \left\lceil \frac{n}{n_{\text{subset}}} \right\rceil \right)$ reduced problems.

4. Numerical Stability of the Solutions $\hat{\mathbf{z}}^{(r)}$ and $\hat{\mathbf{z}}^*$

The numerical stability of solutions in each step of the CLSP estimator — both the (iterated if $r > 1$) first-step estimate $\hat{\mathbf{z}}^{(r)}$ (alternatively, $\hat{\mathbf{z}}_{i_E, j_E}^{(r)} = \left\{ \hat{\mathbf{z}}_{\text{reduced}, 1:m_{\text{subset}}, 1:n_{\text{subset}}}^{(r)}, (i_E, j_E) \right\}$) and the final solution $\hat{\mathbf{z}}^*$ (alternatively, $\hat{\mathbf{z}}_{i_E, j_E}^* = \left\{ \hat{\mathbf{z}}_{\text{reduced}, 1:m_{\text{subset}}, 1:n_{\text{subset}}}^*, (i_E, j_E) \right\}$) — depends on the condition number of the design matrix, $\kappa(\mathbf{A}^{(r)})$, which, given its block-wise formulation in Equation (3), can be analyzed as a function of the “canonical” constraints matrix $\mathbf{C}_{\text{canon}} = [\mathbf{C} \ \mathbf{S}] \in \mathbb{R}^{k \times n}$ as a fixed, r -invariant, part of $\mathbf{A}^{(r)}$. For both full and reduced problems, **sensitivity analysis** of $\kappa(\mathbf{A}^{(r)})$ with respect to constraints (i.e., rows) in $\mathbf{C}_{\text{canon}}$, especially focused on their reduction (i.e., dropping rows), is performed based on the decomposition:

$$\mathbf{A}^{(r)} = \mathbf{B}^{(r)} \mathbf{C}_{\text{canon}} + \mathbf{R}^{(r)} = \mathbf{B}^{(r)} \begin{bmatrix} \mathbf{C} & \mathbf{S} \end{bmatrix} + \mathbf{R}^{(r)}, \quad \mathbf{R}^{(r)} = \mathbf{A}^{(r)} (\mathbf{I} - \mathbf{P}_{\mathcal{R}([\mathbf{C} \ \mathbf{S}]^\top)}) \quad (7)$$

where $\mathbf{B}^{(r)} = \mathbf{A}^{(r)} \mathbf{C}_{\text{canon}}^\dagger \in \mathbb{R}^{m \times k}$, denotes the constraint-induced, r -variant, part of $\mathbf{A}$, and, when the rows of $\mathbf{A}^{(r)}$ lie entirely within $\mathcal{R}(\mathbf{C}_{\text{canon}})$, i.e., $\mathbf{A}^{(r)} = \mathbf{A}^{(r)} \mathbf{P}_{\mathcal{R}([\mathbf{C} \ \mathbf{S}]^\top)} \Leftrightarrow \mathbf{R}^{(r)} = \mathbf{0}$,

$$\begin{aligned} \kappa(\mathbf{A}^{(r)}) &\leq \kappa(\mathbf{B}^{(r)}) \cdot \kappa(\mathbf{C}_{\text{canon}}) \\ \kappa(\mathbf{A}^{(r)}) + \Delta\kappa_A &\leq \left(\kappa(\mathbf{B}^{(r)}) + \Delta\kappa_B \right) \cdot \left(\kappa(\mathbf{C}_{\text{canon}}) + \Delta\kappa_C \right) \\ &= \kappa(\mathbf{B}^{(r)}) \cdot \kappa(\mathbf{C}_{\text{canon}}) + \Delta\kappa_B \cdot \kappa(\mathbf{C}_{\text{canon}}) + \kappa(\mathbf{B}^{(r)}) \cdot \Delta\kappa_C + \Delta\kappa_B \cdot \Delta\kappa_C \end{aligned} \quad (8)$$

i.e., a change in $\kappa(\mathbf{A}^{(r)})$ consists of a (lateral) constraint-side effect, $\kappa(\mathbf{B}^{(r)}) \cdot \Delta\kappa_C$, induced by structural modifications in $\mathbf{C}_{\text{canon}}$, an (axial) model-side effect, $\Delta\kappa_B \cdot \kappa(\mathbf{C}_{\text{canon}})$, reflecting propagation into $\mathbf{B}^{(r)}$ — where $\mathbf{B}^{(r)} = \mathbf{A}^{(r)} \mathbf{C}_{\text{canon}}^\dagger$ and $\mathbf{A}^{(r)} = f(\mathbf{C}_{\text{canon}}, \mathbf{M}_{\text{canon}}^{(r)})$, $\mathbf{M}_{\text{canon}}^{(r)} = g(\mathbf{M}, \varepsilon^{(r-1)})$, and $\varepsilon^{(r-1)} = \mathbf{b} - \mathbf{A}^{(r-1)}\hat{\mathbf{z}}^{(r-1)}$ (for a “partitioned” decomposition of $\mathbf{B}^{(r)}$ given the structure of $\mathbf{C}_{\text{canon}}^\dagger$, see **Appendix A.2**), — and the (non-linear) term, $\Delta\kappa_B \cdot \Delta\kappa_C$.

The condition number of $\mathbf{C}_{\text{canon}} \in \mathbb{R}^{k \times n}$, $\kappa(\mathbf{C}_{\text{canon}})$, the change of which determines the constraint-side effect, $\kappa(\mathbf{B}^{(r)}) \cdot \Delta\kappa_C$, can be itself analyzed with the help of pairwise angular measure between row vectors $\mathbf{c}_i^\top, \mathbf{c}_j^\top \in \mathbf{C}_{\text{canon}}$, where $i, j \in \{1, \dots, k\}$. For any two $\mathbf{c}_i^\top, \mathbf{c}_j^\top \in \mathbb{R}^{1 \times n}$ (alternatively, in a centered form, $\tilde{\mathbf{c}}_i^\top = \mathbf{c}_i^\top - \bar{c}_i$, $\tilde{\mathbf{c}}_j^\top = \mathbf{c}_j^\top - \bar{c}_j \in \mathbb{R}^{1 \times n}$, with $\bar{c}_i = \frac{1}{n} \sum_{k=1}^n c_{i,k}$, $\bar{c}_j = \frac{1}{n} \sum_{k=1}^n c_{j,k}$, $i \neq j$), the cosine of the angle between them $\cos(\theta_{i,j})$ (alternatively, a Pearson correlation coefficient $\rho = \cos(\tilde{\theta}_{i,j})$), captures their angular alignment — values close to ± 1 indicate near-collinearity and potential ill-conditioning, whereas values near 0 imply closeness to orthogonality and improved numerical stability:

$$\forall \mathbf{c}_i^\top, \mathbf{c}_j^\top \in \mathbf{C}_{\text{canon}}, i \neq j \quad \cos(\theta_{i,j}) = \frac{\mathbf{c}_i^\top \cdot \mathbf{c}_j}{\|\mathbf{c}_i^\top\|_2 \cdot \|\mathbf{c}_j^\top\|_2} \in [-1, 1] \quad (9)$$

The pairwise angular measure can be aggregated across j (i.e., for each row $\mathbf{c}_i^\top$) and jointly across i, j (i.e., for the full set of $\binom{k}{2} = \frac{k(k-1)}{2}$ unique row pairs of $\mathbf{C}_{\text{canon}}$), yielding a scalar summary statistic of angular alignment, e.g., the *Root Mean Square Alignment* (RMSA):

$$\text{RMSA}_i = \sqrt{\frac{1}{k-1} \sum_{\substack{j=1 \\ j \neq i}}^k \cos^2(\theta_{i,j})}, \quad \text{RMSA} = \sqrt{\frac{2}{k(k-1)} \sum_{i=1}^{k-1} \sum_{j=i+1}^k \cos^2(\theta_{i,j})} \quad (10)$$

where RMSA_i captures the average angular alignment of constraint i with the rest, while RMSA evaluates the overall constraint anisotropy of $\mathbf{C}_{\text{canon}}$. As with individual cosines $\cos(\theta_{i,j})$ (alternatively, $\rho = \cos(\tilde{\theta}_{i,j})$), values near 1 suggest collinearity and potential numerical instability, while values near 0 reflect angular dispersion and reduced condition number $\kappa(\mathbf{C}_{\text{canon}})$. The use of squared cosines $\cos^2(\theta_{i,j})$ (alternatively, $\rho^2 = \cos^2(\tilde{\theta}_{i,j})$) in contrast with other metrics (e.g., absolute values) maintains consistency with the Frobenius norms in $\kappa(\mathbf{C}_{\text{canon}})$ and assigns greater penalization to near-collinear constraints in $\mathbf{C}_{\text{canon}}$. Furthermore, a combination of RMSA_i , $i \in \{1, \dots, k\}$, and the above-described *marginal effects*, obtained by excluding row i from $\mathbf{C}_{\text{canon}}$ — namely, $\kappa(\mathbf{B}^{(r)}) \cdot \Delta\kappa_C^{(-i)}$, $\Delta\kappa_B^{(-i)} \cdot \kappa(\mathbf{C}_{\text{canon}})$, and $\Delta\kappa_B^{(-i)} \cdot \Delta\kappa_C^{(-i)}$ if $\mathbf{R}^{(r)} = \mathbf{0}$ — together with the corresponding changes in the *goodness-of-fit statistics* from the solutions $\hat{\mathbf{z}}^{(r)(-i)}$ and $\hat{\mathbf{z}}^{*(-i)}$ (consult Section 5), can be visually presented over either all i , or a filtered subset, such that $\text{RMSA}_i \geq \text{threshold}$, as depicted in Figure 5.

Constraint	RMSA_i (total: ω)	$\Delta\kappa_C^{(-i)}$	$\Delta\kappa_B^{(-i)}$	$\Delta\kappa_A^{(-i)}$	$ \text{GoF}^{(-i)} $	
##	-----	$\pm\lambda$	$\pm\lambda$	$\pm\lambda$	$\pm\lambda$	
##	-----	$\pm\lambda$	$\pm\lambda$	$\pm\lambda$	$\pm\lambda$	
##	-----	$\pm\lambda$	$\pm\lambda$	$\pm\lambda$	$\pm\lambda$	
##	-----	$\pm\lambda$	$\pm\lambda$	$\pm\lambda$	$\pm\lambda$	
##	-----	$\pm\lambda$	$\pm\lambda$	$\pm\lambda$	$\pm\lambda$	
##	-----	$\pm\lambda$	$\pm\lambda$	$\pm\lambda$	$\pm\lambda$	
##	-----	$\pm\lambda$	$\pm\lambda$	$\pm\lambda$	$\pm\lambda$	

Figure 5. Schematic correlogram of constraint rows i with high alignment ($\text{RMSA}_i \geq \text{threshold}$). Each row encodes directional alignment and symbolic marginal effects on model conditioning and fit.

5. Goodness of Fit of the Solutions $\hat{\mathbf{z}}^{(r)}$ and $\hat{\mathbf{z}}^*$

Given the hybrid structure of the CLSP estimator — comprising a least-squares-based (iterated if $r > 1$) initial estimate $\hat{\mathbf{z}}^{(r)}$ (or $\hat{\mathbf{z}}_{i_E, j_E}^{(r)} = \left\{ \hat{\mathbf{z}}_{\text{reduced}, 1:m_{\text{subset}}, 1:n_{\text{subset}}}^{(r), (i_E, j_E)} \right\}$), followed by a regularization-based final solution $\hat{\mathbf{z}}^*$ (or $\hat{\mathbf{z}}_{i_E, j_E}^* = \left\{ \hat{\mathbf{z}}_{\text{reduced}, 1:m_{\text{subset}}, 1:n_{\text{subset}}}^{*, (i_E, j_E)} \right\}$) — and the potentially ill-posed or underdetermined nature of the problems it addresses, standard **goodness-of-fit methodology** — such as explained variance measures (e.g., the coefficient of determination, R^2), error-based metrics (e.g., Root Mean Square Error, RMSE, for comparability with the above-described RMSA), hypothesis tests (e.g., F -tests

in the analysis of variance and t -tests at the coefficient level), or confidence intervals (e.g., based on Student's t - and the normal distribution) — requires appropriate adaptation to be uniformly applied to both full and reduced problems, with some statistics applicable only to Step 1 and/or under specific conditions (namely, the ANOVA F - tests, t -tests, and confidence intervals). In response, several more robust metrics are developed below: partial R^2 , normalized RMSE, t -test for the mean of the NRMSE, and a diagnostic interval based on the condition number.

For the *explained variance measures*, the block-wise formulation of the full vector of unknowns $\mathbf{z}^{(r)} = \begin{bmatrix} \mathbf{x}^{(r)} \\ (\mathbf{y}^{(r)})^* \end{bmatrix} \in \mathbb{R}^n$, — obtained from $\mathbf{x}^{(r)} = \mathbf{Q}_X \begin{bmatrix} \mathbf{x}_M^{(r)} \\ \mathbf{x}_L^{(r)} \end{bmatrix} \in \mathbb{R}^p$ (where $\mathbf{x}_M^{(r)} \in \mathbb{R}^{p-l}$ is the vector of model (target) variables and $\mathbf{x}_L^{(r)} \in \mathbb{R}^l$, $l \geq 0$, is a vector of latent variables present only in the constraints) and a vector of slack variables $(\mathbf{y}^{(r)})^* \in \mathbb{R}^{n-p}$, — the design matrix $\mathbf{A}^{(r)} = \begin{bmatrix} \mathbf{C}_{\text{canon}} \\ \mathbf{M}_{\text{canon}}^{(r)} \end{bmatrix} \in \mathbb{R}^{m \times n}$, where $\mathbf{M}_{\text{canon}}^{(r)} = [\mathbf{M} \ \mathbf{0}, r = 1; \mathbf{S}_{\text{residual}}, r > 1] \in \mathbb{R}^{(m-k) \times n}$ is the “canonical” model matrix, and the input $\mathbf{b} = \begin{bmatrix} \mathbf{b}_C \\ \mathbf{b}_M \end{bmatrix} \in \mathbb{R}^m$, all necessitate a “partial” R^2 to isolate the estimated variables $\hat{\mathbf{x}}^*$, where, given the vector $\mathbf{1} = \{1, \dots, 1\}^\top$,

$$R_{\text{partial}}^2 = 1 - \frac{\|\mathbf{b}_M - \mathbf{M}\hat{\mathbf{x}}^*\|_2^2}{\|\mathbf{b}_M - \bar{b}_M \mathbf{1}\|_2^2}, \quad \bar{b}_M = \frac{1}{m-k} \sum_{i=1}^{m-k} b_{M,i} \quad (11)$$

provided that $\text{rank}(\mathbf{A}^{(r)}) = n$, $m \geq n$, i.e., applicable strictly to (over)determined problems. If $\mathbf{C} = \emptyset \Leftrightarrow k = 0 \wedge \mathbf{y} = \emptyset \Leftrightarrow n = p$, then $\mathbf{A}^{(r)} \equiv \mathbf{M}$ and R_{partial}^2 reduces to the classical R^2 .

Next, for the *error-based metrics*, independent of $\text{rank}(\mathbf{A}^{(r)})$ and thus more robust across underdetermined cases (in contrast to R_{partial}^2), a Root Mean Square Error (RMSE) can be defined for both (a) the full solution $\hat{\mathbf{z}}^*$, via $\text{RMSE} = \frac{1}{\sqrt{m}} \|\mathbf{b} - \mathbf{A}^{(r)} \hat{\mathbf{z}}^*\|_2$, serving as an overall measure of fit, and (b) its variable component $\hat{\mathbf{x}}^*$, via $\text{RMSE}_{\text{partial}}^{(r)} = \frac{1}{\sqrt{m-k}} \|\mathbf{b}_{k+1:m} - \mathbf{M}\hat{\mathbf{x}}^*\|_2$, quantifying the fit quality with respect to the estimated target variables only (in alignment with R_{partial}^2). Then, to enable cross-model comparisons and especially its use in hypothesis testing (see below), RMSE must be normalized — typically by the standard deviation of the reference input, $\sigma_{\mathbf{b}}$ or $\sigma_{\mathbf{b}_M}$, — yielding the **Normalized Root Mean Square Error (NRMSE)** comparable across datasets and/or models (e.g., in the t -test for the mean of the NRMSE):

$$\text{NRMSE} = \frac{\frac{1}{\sqrt{m}} \|\mathbf{b} - \mathbf{A}^{(r)} \hat{\mathbf{z}}^*\|_2}{\sigma_{\mathbf{b}}}, \quad \text{NRMSE}_{\text{partial}} = \frac{\frac{1}{\sqrt{m-k}} \|\mathbf{b}_{k+1:m} - \mathbf{M}\hat{\mathbf{x}}^*\|_2}{\sigma_{\mathbf{b}_M}} \quad (12)$$

The standard deviation is preferred in statistical methodology over the alternatives (e.g., max-min scaling or range-based metrics) because it expresses residuals in units of their variability, producing a scale-free measure analogous to the standard normal distribution. It is also more common in practical implementations (such as Python's `sklearn` or MATLAB).

Also, for the *hypothesis tests*, classical ANOVA-based F -tests and coefficient-level t -tests — relying on variance decomposition and residual degrees of freedom — are applicable exclusively to the least-squares-based (iterated if $r > 1$) first-step estimate $\hat{\mathbf{z}}^{(r)}$ and to the final solution $\hat{\mathbf{z}}^* = \hat{\mathbf{z}}^{(r)}$, given a strictly overdetermined $\mathbf{A}^{(r)} \mathbf{z}^{(r)} = \mathbf{b}$, i.e., $\text{rank}(\mathbf{A}^{(r)}) = n$, $m > n$ (an OLS case), and under the assumption of homoscedastic and normally distributed residuals. Then, in the case of the F -tests, the test statistics follow the distributions (a) $F \sim \mathcal{F}(q, m-n)$ and (b) $F_{\text{partial}} \sim \mathcal{F}(q_{\text{partial}}, m-n)$, a Wald-type test on a linear restriction — with the degrees of freedom $q = n-1$ and $q_{\text{partial}} = p-1$ if $\mathbf{x}^{(r)}$ is a vector of coefficients from a linear(ized) model with an intercept and $q = n$ and $q_{\text{partial}} = p$ otherwise, — yielding (a) $H_0 : \hat{\mathbf{z}}^{(r)} = \mathbf{0}$ and (b) $H_0^{\text{partial}} : \hat{\mathbf{x}}^{(r)} = \mathbf{R}\hat{\mathbf{z}}^{(r)} = \mathbf{0}$, where $\mathbf{R} \in \mathbb{R}^{p \times n}$ is a restriction matrix:

$$F = \frac{\frac{1}{q} \|\mathbf{A}^{(r)} \hat{\mathbf{z}}^{(r)}\|_2^2}{\frac{1}{m-n} \|\mathbf{b} - \mathbf{A}^{(r)} \hat{\mathbf{z}}^{(r)}\|_2^2}, \quad F_{\text{partial}} = \frac{\frac{1}{q_{\text{partial}}} (\mathbf{R}\hat{\mathbf{z}}^{(r)})^\top \left[\mathbf{R}((\mathbf{A}^{(r)})^\top \mathbf{A}^{(r)})^{-1} \mathbf{R}^\top \right]^{-1} (\mathbf{R}\hat{\mathbf{z}}^{(r)})}{\frac{1}{m-n} \|\mathbf{b} - \mathbf{A}^{(r)} \hat{\mathbf{z}}^{(r)}\|_2^2} \quad (13)$$

In the case of the t -tests, (a) $\forall_{i \in \{1, \dots, n\}} t_i \sim t_{m-n}$ (n test statistics) and (b) $\forall_{i \in \{1, \dots, p\}} t_{\text{partial}, i} \equiv t_i$ (p test statistics), given that the partial result $\hat{\mathbf{x}}^{(r)}$ or $\hat{\mathbf{x}}^* = \hat{\mathbf{x}}^{(r)}$ is a subvector of length p of $\hat{\mathbf{z}}^{(r)}$ (also $\hat{\mathbf{z}}^* = \hat{\mathbf{z}}^{(r)}$), yielding (a) $\forall_{i \in \{1, \dots, n\}} H_0 : \hat{z}_i^{(r)} = 0$ and (b) $\forall_{i \in \{1, \dots, p\}} H_0^{\text{partial}} : \hat{x}_i^{(r)} = 0$:

$$\forall_{i \in \{1, \dots, n\}} t_i = \frac{\hat{z}_i^{(r)}}{\left(\frac{1}{\sqrt{m-n}} \|\mathbf{b} - \mathbf{A}^{(r)} \hat{\mathbf{z}}^{(r)}\|_2 \right) \cdot \sqrt{[(\mathbf{A}^{(r)})^\top \mathbf{A}^{(r)}]^{-1}_{ii}}}, \quad \forall_{i \in \{1, \dots, p\}} t_{\text{partial}, i} \equiv t_i \quad (14)$$

An alternative hypothesis test — robust to both $\text{rank}(\mathbf{A}^{(r)})$ and the step of the CLSP algorithm, i.e., valid for the final solution $\hat{\mathbf{z}}^*$ (in contrast to the classical ANOVA-based F -tests or coefficient-level t -tests) — can be a **t -test for the mean of the NRMSE**, comparing the observed $\text{NRMSE}_{\text{obs}}$ (as μ_0), both full and partial, to the mean of a simulated sample $\text{NRMSE}_{\text{sim}}^{(\tau)}$, generated via *Monte Carlo simulation*, typically from a uniformly (a structureless baseline) or normally (a “canonical” choice) distributed random input $\mathbf{b}^{(\tau)} \sim \mathcal{U}(\mathbf{0}, \mathbf{I}_m)$ or $\mathbf{b}^{(\tau)} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_m)$ — both the distributions being standard, — for $\tau = 1, \dots, T$, where T is the sample size, yielding $t \sim t_{T-1}$ with $H_0 : \text{NRMSE}_{\text{obs}} = \overline{\text{NRMSE}}_{\text{sim}}$ and $H_1 : \text{NRMSE}_{\text{obs}} \neq \overline{\text{NRMSE}}_{\text{sim}}$ for the two-sided test and $H_1 : \text{NRMSE}_{\text{obs}} \geq \overline{\text{NRMSE}}_{\text{sim}}$ for the one-sided one:

$$t = \frac{\overline{\text{NRMSE}}_{\text{sim}} - \text{NRMSE}_{\text{obs}}}{\sqrt{\frac{1}{T(T-1)} \sum_{\tau=1}^T \left(\text{NRMSE}_{\text{sim}}^{(\tau)} - \overline{\text{NRMSE}}_{\text{sim}} \right)^2}}, \quad \overline{\text{NRMSE}}_{\text{sim}} = \frac{1}{T} \sum_{\tau=1}^T \text{NRMSE}_{\text{sim}}^{(\tau)} \quad (15)$$

justified when $\mathbf{b}^{(\tau)}$ is normally distributed or approximately normally distributed based on the Lindeberg-Lévy Central Limit Theorem (CLT), i.e., when $T > 30$ (as a rule of thumb) for $\mathbf{b}^{(\tau)} \sim \mathcal{U}(\mathbf{0}, \mathbf{I}_m)$. Then, H_0 denotes a good fit for $\hat{\mathbf{z}}^*$ in the sense that $\text{NRMSE}_{\text{obs}}$ does not significantly deviate from the simulated distribution, i.e., H_0 should not be rejected for the CLSP model to be considered statistically consistent. In practical implementations (such as Python, R, Stata, SAS, Julia, and MATLAB/Octave), T typically ranges from 50 to 50,000.

Finally, for the *confidence intervals*, classical formulations for (a) $z_i^{(r)}$, $i \in \{1, \dots, n\}$, and (b) $x_i^{(r)}$, $i \in \{1, \dots, p\}$, — based on Student’s t - and, asymptotically, the standard normal distribution, $t \sim t_{m-n}$ and $z \sim \mathcal{N}(0, 1)$, — are equivalently exclusively applicable to the least-squares-based (iterated if $r > 1$) first-step estimate $\hat{\mathbf{z}}^{(r)}$ and to the final solution $\hat{\mathbf{z}}^* = \hat{\mathbf{z}}^{(r)}$ under a strictly overdetermined $\mathbf{A}^{(r)} \mathbf{z}^{(r)} = \mathbf{b}$, i.e., $\text{rank}(\mathbf{A}^{(r)}) = n$, $m > n$ (an OLS case), and assuming homoscedastic and normally distributed residuals. Then, provided $\alpha \in (0, 1)$, such as $\alpha \in \{0.01, 0.05, 0.1\}$, the confidence intervals for (a) $z_i^{(r)}$ and (b) $x_i^{(r)}$ are:

$$\begin{aligned} \forall_{i \in \{1, \dots, n\}} z_i^{(r)} &\in \hat{z}_i^{(r)} \pm t_{m-n, 1-\alpha/2} \cdot \frac{1}{\sqrt{m-n}} \|\mathbf{b} - \mathbf{A}^{(r)} \hat{\mathbf{z}}^{(r)}\|_2 \cdot \sqrt{[(\mathbf{A}^{(r)})^\top \mathbf{A}^{(r)}]^{-1}_{ii}}, \\ \forall_{i \in \{1, \dots, p\}} x_i^{(r)} &\equiv z_i^{(r)} \end{aligned} \quad (16)$$

given that the partial result $\hat{\mathbf{x}}^{(r)}$ or $\hat{\mathbf{x}}^* = \hat{\mathbf{x}}^{(r)}$ is a subvector of length p of $\hat{\mathbf{z}}^{(r)}$ or $\hat{\mathbf{z}}^* = \hat{\mathbf{z}}^{(r)}$. For $n > 30$, the distribution of the test statistic, $t_{m-n, 1-\alpha/2}$, approaches a standard normal one, $Z_{1-\alpha/2}$, based on the CLT, yielding an alternative notation for the confidence intervals:

$$\begin{aligned} \forall_{i \in \{1, \dots, n\}} z_i^{(r)} &\in \hat{z}_i^{(r)} \pm Z_{1-\alpha/2} \cdot \frac{1}{\sqrt{m-n}} \|\mathbf{b} - \mathbf{A}^{(r)} \hat{\mathbf{z}}^{(r)}\|_2 \cdot \sqrt{[(\mathbf{A}^{(r)})^\top \mathbf{A}^{(r)}]^{-1}_{ii}}, \\ \forall_{i \in \{1, \dots, p\}} x_i^{(r)} &\equiv z_i^{(r)} \end{aligned} \quad (17)$$

An alternative “confidence interval” — equivalently robust to both $\text{rank}(\mathbf{A}^{(r)})$ and the step of the CLSP algorithm, i.e., valid for the final solution $\hat{\mathbf{z}}^*$ (in contrast to the classical Student’s t - and, asymptotically, the standard normal distribution-based formulations) — can be *constructed*

deterministically via condition-weighted bands, relying on componentwise numerical sensitivity. Let the residual vector $\mathbf{b} - \mathbf{A}^{(r)}\hat{\mathbf{z}}^*$ be a (backward) perturbation in $\mathbf{b}$, i.e., $\Delta\mathbf{b} = \mathbf{b} - \mathbf{A}^{(r)}\hat{\mathbf{z}}^*$. Then, squaring both sides of the classical first-order inequality $\frac{\|\Delta\hat{\mathbf{z}}^*\|_2}{\|\hat{\mathbf{z}}^*\|_2} \leq \kappa(\mathbf{A}^{(r)}) \cdot \frac{\|\Delta\mathbf{b}\|_2}{\|\mathbf{b}\|_2}$ yields $\frac{\sum_{i=1}^n (\Delta\hat{z}_i^*)^2}{\sum_{i=1}^n (\hat{z}_i^*)^2} \leq \kappa(\mathbf{A}^{(r)})^2 \cdot \frac{\|\mathbf{b} - \mathbf{A}^{(r)}\hat{\mathbf{z}}^*\|_2^2}{\|\mathbf{b}\|_2^2}$, where $\kappa(\mathbf{A}^{(r)}) = \|\mathbf{A}^{(r)}\|_2 \|(\mathbf{A}^{(r)})^\dagger\|_2 = \frac{\sigma_{\max}(\mathbf{A}^{(r)})}{\sigma_{\min}(\mathbf{A}^{(r)})}$ — $\sigma_{\max}$ and $\sigma_{\min}$ being the biggest and smallest singular values of $\mathbf{A}^{(r)}$ — iff $\sigma_{\min}(\mathbf{A}^{(r)}) > 0 \wedge \|\mathbf{b}\|_2 > 0$. Under a uniform relative squared perturbation (a heuristic allocation of error across components as a simplification, not an assumption) of $\mathbf{z}^*$, rearranging terms and taking square roots of both sides gives $\forall_{i \in \{1, \dots, n\}} |\Delta\hat{z}_i^*| \leq |\hat{z}_i^*| \cdot \kappa(\mathbf{A}^{(r)}) \cdot \frac{\|\mathbf{b} - \mathbf{A}^{(r)}\hat{\mathbf{z}}^*\|_2}{\|\mathbf{b}\|_2}$ and a condition-weighted “confidence” band:

$$\forall_{i \in \{1, \dots, n\}} z_i^* \in \hat{z}_i^* \cdot \left(1 \pm \kappa(\mathbf{A}^{(r)}) \cdot \frac{\|\mathbf{b} - \mathbf{A}^{(r)}\hat{\mathbf{z}}^*\|_2}{\|\mathbf{b}\|_2} \right), \quad \forall_{i \in \{1, \dots, p\}} x_i^* \equiv z_i^* \quad (18)$$

which is a **diagnostic interval based on the condition number** for $\mathbf{z}^*$, consisting of (1) a canonical-form system conditioning, $\kappa(\mathbf{A}^{(r)})$, (2) a normalized model misfit, $\frac{\|\mathbf{b} - \mathbf{A}^{(r)}\hat{\mathbf{z}}^*\|_2}{\|\mathbf{b}\|_2}$, and (3) the “scale” of the final solution, $\hat{\mathbf{z}}^*$, without probabilistic assumptions. A perturbation in one or more of these three components, e.g., caused by a change in RMSA resulting from dropping selected constraints (consult Section 4), will affect the limits of $\mathbf{z}^*$. Under a perfect fit, the interval collapses to $z_i^* \in \hat{z}_i^* \pm 0$ and, for $\kappa(\mathbf{A}^{(r)}) \gg 1$, tends to be very wide.

To summarize, the CLSP estimator requires a goodness-of-fit framework that reflects both its algorithmic structure and numerical stability. Classical methods remain informative but are valid only under strict overdetermination ($m > n$), full-rank design matrix $\mathbf{A}^{(r)}$, and distributional assumptions (primarily, homoscedasticity and normality of residuals). In contrast, the proposed alternatives — (a) partial R^2_{partial} , (b) NRMSE and NRMSE_{partial}, (c) t -tests for the mean of NRMSE with the help of Monte Carlo sampling, and (d) condition-weighted confidence bands — are robust to rank deficiency and ill-posedness, making them applicable in practical implementations (namely, CLSP packages in Python, R, and Stata).

6. Special Cases of CLSP Problems: APs, CMLS, and LPRLS/QPRLS

The structure of the *design matrix* $\mathbf{A}^{(r)}$ in the CLSP estimator, as defined in Equation (3) (consult Section 3), allows for accommodating a selection of special-case problems, out of which three cases are covered, but the list may not be exhaustive. Among others, the CLSP estimator is, given its ability to address ill-posed or underdetermined problems under linear(ized) constraints, efficient in addressing what can be referred to as **allocation problems (APs)** (or, for flow variables, **tabular matrix problems, TMs**) — i.e., in most cases, underdetermined problems involving matrices of dimensions $m \times p$ to be estimated, subject to known row and column sums, with the degrees of freedom (i.e., nullity) equal to $n - \text{rank}(\mathbf{A}^{(r)}) \geq mp + s^* - (\frac{m}{i} + \frac{p}{j}) - q$, where $0 \leq s^* \leq (\frac{m}{i} + \frac{p}{j}) + q$ is the number of active (non-zero) slack variables and $0 \leq q \leq \text{rank}(\mathbf{M}_{\text{partial}}) \leq \min(k, mp)$ is the number of known model (target) variables (e.g., a zero diagonal), — whose design matrix, $\mathbf{A}^{(r)} \in \mathbb{R}^{(\frac{m}{i} + \frac{p}{j} + k) \times n}$, comprises a (by convention) *row-sum-column-sum constraints matrix* $\mathbf{C} = \left[\left((\mathbf{I}_m \otimes \mathbf{1}_i) \otimes \mathbf{1}_p \right)^\top \left(\mathbf{1}_m \otimes \left(\mathbf{I}_p \otimes \mathbf{1}_j \right) \right)^\top \right]^\top \in \mathbb{R}^{(\frac{m}{i} + \frac{p}{j}) \times mp}$ (where $\mathbf{1} = \{1, \dots, 1\}$), a *model matrix* $\mathbf{M} = \mathbf{Q}_M^L \begin{bmatrix} \mathbf{M}_{\text{partial}} \\ \mathbf{0} \end{bmatrix} \mathbf{Q}_M^R \in \mathbb{R}^{k \times mp}$ — in a trivial case, $\mathbf{M}_{\text{partial}} \subseteq \mathbf{I}_{mp}$ — a *sign slack matrix* $\mathbf{S} = \mathbf{Q}_S^L \begin{bmatrix} \mathbf{S}_s \subseteq \forall_{\sigma_1, \dots, \sigma_s \in \{-1, +1\}} \text{diag}(\sigma_1, \dots, \sigma_s) \\ \mathbf{0} \end{bmatrix} \mathbf{Q}_S^R \in \mathbb{R}^{(\frac{m}{i} + \frac{p}{j}) \times (n - mp)}$, and

either a zero matrix $\mathbf{0}$ (if $r = 1$) or a (standard) reverse-sign slack matrix $\mathbf{S}_{\text{residual}} \in \mathbb{R}^{k \times (n-mp)}$ (provided $r > 1$):

$$\mathbf{A}^{(r)} = \begin{bmatrix} \begin{bmatrix} (\mathbf{I}_m^i \otimes \mathbf{1}_i) \otimes \mathbf{1}_p \\ \mathbf{1}_m \otimes (\mathbf{I}_p^j \otimes \mathbf{1}_j) \end{bmatrix} & \mathbf{S} \\ \mathbf{M} & \begin{cases} \mathbf{0}, & \text{initial iteration, } r = 1 \\ \mathbf{S}_{\text{residual}}, & \text{subsequent iterations, } r > 1 \end{cases} \end{bmatrix} \quad (19)$$

where $\frac{m}{i}$ and $\frac{p}{j}$ denote groupings of the m rows and p columns, respectively, into i and j homogeneous blocks (when no grouped row or column sums are available, $i = j = 1$); with real-world examples including: (a) input-output tables (national, inter-country, or satellite); (b) structural matrices (e.g., trade, country-product, investment, or migration); (c) financial clearing and budget-balancing; and (d) data interpolations (e.g., quarterly data from annual totals). Given the available literature in 2025, the first pseudoinverse-based method of estimating input-output tables was proposed in Pereira-López et al. [95] and the first APs (TMs)-based study was conducted in Bolotov [16], attempting to interpolate a world-level (in total, 232 countries and dependent territories) “bilateral flow” matrix of foreign direct investment (FDI) for the year 2013, based on known row and column totals from UNCTAD data (aggregate outward and inward FDI). The estimation employed a minimum-norm least squares solution under APs (TMs)-style equality constraints, based on the Moore-Penrose pseudoinverse of a “generalized Leontief structural matrix”, rendering it equivalent to the initial (non-regularized) solution in the CLSP framework, $\hat{\mathbf{z}}^{(1)} = \mathbf{A}_Z^{(1)\dagger} \mathbf{b}$. As a rule of thumb, AP (TM)-type problems are rank deficient for $m > 2 \wedge p > 2$ with the CLSP being a *unique* (if $\alpha \in (0, 1]$) and an *MNBLUE* (if $\alpha = 1$) estimator (consult Section 3).

In addition to the APs (TMs), the CLSP estimator is also efficient in addressing what can be referred to as **constrained-model least squares (CMLS)** (or, more generally, **regression problems, RPs**) — i.e., (over)determined problems involving vectors of dimension p to be estimated with (standard) OLS degrees of freedom, subject to, among others, linear(ized) data matrix $\mathbf{D} \in \mathbb{R}^{k \times p}$ -related (in)equality constraints, $\forall_{i \in \{1, \dots, t\}} \mathbf{C}_i \equiv \mathbf{T}_i \mathbf{D} \in \mathbb{R}^{m_i \times p}$, where $\mathbf{T}_i \in \mathbb{R}^{m_i \times k}$ is a transformation matrix of $\mathbf{D}$, such as the i -th difference or shift (lag/lead), and $\mathbf{C}_i \mathbf{x} \geq \gamma_i$, — whose design matrix, $\mathbf{A}^{(r)} \in \mathbb{R}^{(\sum_{i=1}^u m_i + k) \times n}$, consists of a u -block constraints matrix $\mathbf{C} = \begin{bmatrix} \mathbf{C}_1 \\ \vdots \\ \mathbf{C}_u \end{bmatrix} \in \mathbb{R}^{(\sum_{i=1}^u m_i) \times p}$ ($\forall_{i \in \{1, \dots, u\}} \mathbf{C}_i \in \mathbb{R}^{m_i \times p}$), $u \geq t$, a data matrix as the *model matrix* $\mathbf{M} \equiv \mathbf{D} \in \mathbb{R}^{k \times p}$, a (standard) sign slack matrix $\mathbf{S} \in \mathbb{R}^{(\sum_{i=1}^u m_i) \times (n-p)}$, and either a zero matrix $\mathbf{0}$ (if $r = 1$) or a (standard) reverse-sign slack matrix $\mathbf{S}_{\text{residual}} \in \mathbb{R}^{k \times (n-p)}$ (when $r > 1$):

$$\mathbf{A}^{(r)} = \begin{bmatrix} \begin{bmatrix} \mathbf{C}_1 \\ \vdots \\ \mathbf{C}_u \end{bmatrix} & \mathbf{S} \\ \mathbf{D} & \begin{cases} \mathbf{0}, & \text{initial iteration, } r = 1 \\ \mathbf{S}_{\text{residual}}, & \text{subsequent iterations, } r > 1 \end{cases} \end{bmatrix} \quad (20)$$

with real-world examples including: (a) previously infeasible textbook-definition econometric models of economic (both micro and macro) variables (e.g., business-cycle models), (b) additional constraints applied to classical econometric models (e.g., demand analysis). Given the available literature in 2025, the first studies structurally resembling CMLS (RPs) were conducted in Bolotov [15,96], focusing on the decomposition of the long U.S. GDP time series into trend (y_t) and cyclical (c_t) components — using exponential trend, moving average, Hodrick-Prescott filter, and Baxter-King filter — under constraints on its first difference, based on the U.S. National Bureau of Economic Research (NBER) delimitation of the U.S. business cycle. c_t was then modeled with the help of linear regression (OLS), based on an n -th order polynomial with extreme values smoothed via a factor of $\frac{1}{10^{n-i}}$ and simultaneously penalized via an n -th or $(n+1)$ -th root, $c_t = \beta_0 + \beta_1 \int \left(\frac{n, n+1}{\prod_{i=1}^n (t - t_i)} \cdot \frac{1}{10^{n-i}} \right) dt + \varepsilon_t$, where β_0, β_1 are the model parameters, $\forall_{1 \leq i \leq n} t_i$ are the values of the (externally given) stationary points, and ε_t is the error.

The order of the polynomial and the ad-hoc selection of smoothing and penalizing factors, however, render such a method inferior to the CLSP. Namely, the *unique* (if $\alpha \in (0, 1]$) and an *MNBLUE* (if $\alpha = 1$) CLSP estimator (consult Section 3) allows (a) presence of inequality constraints and (b) their ill-posed formulations.

Finally, the CLSP estimator can be used to address (often unsolvable using classical solvers) underdetermined and/or ill-posed — caused by too few and/or mutually inconsistent or infeasible constraints, as in the sense of Whiteside et al. [93] and Blair [76,94], — LP and QP problems, an approach hereafter referred to as **linear/quadratic programming via regularized least squares (LPRLS/QPRLS)**: CLSP substitutes the original objective function of the LP/QP problem with the canonical form $\mathbf{A}^{(r)}\mathbf{z}^{(r)} = \mathbf{b}$ (i.e., focusing solely on the problem's constraints, without distinguishing between the LP and QP cases) with $\hat{\mathbf{x}}_M^* = \text{Proj}_{\mathbf{x}_M^{(r)}} \left\{ \mathbf{Q}_X \begin{bmatrix} \mathbf{x}_M^{(r)} \\ \mathbf{x}_L^{(r)} \end{bmatrix} \in \mathbb{R}^p : \forall_{1 \leq i \leq s} g_i(\mathbf{x}_M^{(r)}, \mathbf{x}_L^{(r)}) \gtrless \gamma_i, \forall_j h_j(\mathbf{x}_M^{(r)}, \mathbf{x}_L^{(r)}) = \eta_j \right\} \subseteq \mathbf{z}^*$ being the solution of the original problem — where $\mathbf{Q}_X$ is a permutation matrix, and $g_i(\cdot) \gtrless \gamma_i$ and $h_j(\cdot) = \eta_j$ represent the linear(ized) inequality and equality constraints, $\forall_{i,j} \gamma_i, \eta_j \in \mathbf{b}$ — and the degrees of freedom being equal — under the classical formalization of a primal LP/QP problem (consult Section 2) — to $n - \text{rank}(\mathbf{A}^{(r)}) \geq (p + s) - (m_{ub} + m_{eq} + m_{nn}) = p - m_{eq}$, where p is the length of $\mathbf{x}_M^{(r)}$, $s = m_{ub} + m_{nn}$ is the number of introduced slack variables, while $m_{ub} \geq 0$, $m_{eq} \geq 0$, and $m_{nn} \equiv p$ denote the numbers of upper-bound, equality, and non-negativity constraints, respectively (under the standard assumption that all model variables are constrained to be non-negative). In the LPRLS/QPRLS, the design matrix, $\mathbf{A}^{(r)} \in \mathbb{R}^{(m_{ub}+m_{eq}+m_{nn}) \times n}$, by the definition of LP/QP, is r -invariant and consists of a block constraints matrix $\mathbf{C} = \begin{bmatrix} \mathbf{C}_{ub} \\ \mathbf{C}_{eq} \\ \mathbf{C}_{nn} \end{bmatrix} \in \mathbb{R}^{(m_{ub}+m_{eq}+m_{nn}) \times p}$, where $\mathbf{C}_{ub} \in \mathbb{R}^{m_{ub} \times p}$, $\mathbf{C}_{eq} \in \mathbb{R}^{m_{eq} \times p}$, and $\mathbf{C}_{nn} \equiv \mathbf{I}_p \in \mathbb{R}^{p \times p}$, a (standard) sign *slack matrix* $\mathbf{S} \in \mathbb{R}^{(m_{ub}+m_{eq}+m_{nn}) \times (n-p)}$, and $\mathbf{M}_{\text{canon}} = \emptyset$:

$$\mathbf{A}^{(r)} = \mathbf{A}^{(1)} = \begin{bmatrix} \begin{bmatrix} \mathbf{C}_{ub} \\ \mathbf{C}_{eq} \\ \mathbf{C}_{nn} \end{bmatrix} & \mathbf{S} \end{bmatrix} \quad (21)$$

Given the available literature in 2025, the first documented attempts to incorporate LS into LP and QP included, among others, Dax's LS-based steepest-descent algorithm [85] and the primal-dual NNLS-based algorithm (LSPD) [86], whose structural restrictions — in terms of well-posedness and admissible constraints — can be relaxed through LPRLS/QPRLS, still guaranteeing a *unique* (if $\alpha \in (0, 1]$) and an *MNBLUE* (if $\alpha = 1$) solution (consult Section 3).

7. Monte Carlo Experiment and Numerical Examples

The Monte Carlo experiment and the simulated examples presented below illustrate the applicability of the CLSP to the AP (TM), CMLS (RP), and LPRLS/QPRLS problems. The standardized form of $\mathbf{C}$ in the APs (TMs), i.e., $\mathbf{C} = f(m, i, p, j)$, with given $\mathbf{S}$ and r , enables **large-scale simulation of the distribution of NRMSE** — a robust goodness-of-fit metric for the CLSP — through *repeated random trials* under varying matrix dimensions $m \times p$, row and column group sizes i and j , and composition of $\mathbf{M}$, inferring asymptotic percentile means and standard deviations for practical implementations (e.g., the t -test for the mean). This text presents the results of a simulation study on the distribution of the NRMSE from the first-step estimate $\hat{\mathbf{z}}^{(1)}$ (implementable in standard econometric software without CLSP modules) for $i = 1, j = 1, \mathbf{S} = \emptyset, \mathbf{Z} = \mathbf{I}_{m+p+k} \Leftrightarrow (\mathbf{A}_Z^{(1)})^+ = (\mathbf{A}^{(1)})^+$, and $r = 1$, conducted in Stata/MP (set to version 14.0) with the help of Mata's 128-bit floating-point cross-product-based `quadcross()` for greater precision and SVD-based `svsolve()` with a strict tolerance of `c("epsdouble")` for increased numerical stability, and a random-variate seed 123456789 (see **Appendix B.1** containing a list of dependencies [97] and the Stata .do-file). In this 50,000-iterations Monte Carlo simulation study, spanning matrix dimensions from 1×2 and 2×1 up to 50×50 , random normal input vectors $\mathbf{b} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_{m+p+k})$ were generated for each run, applied with and without zero diagonals (i.e., if, under $l = 0 \Leftrightarrow \mathbf{x}_L = \emptyset, \mathbf{x}_M = \mathbf{x} = \mathbf{z}_{1:mp} \in \mathbb{R}^{mp}$ is reshaped into a $(m \times p)$ -matrix $\mathbf{X}_M = \mathbf{X} \in \mathbb{R}^{m \times p}$,

then $\forall_{i \in \{1, \dots, \min(m, p)\}} \mathbf{X}_{ii} = 0$). Thus, a total of 249.9 million observations $(= (50 \cdot 50 - 1) \cdot 2 \cdot 50,000)$ was obtained via the formula $\text{NRMSE} = \frac{1}{\sqrt{m}} \|\mathbf{b} - \mathbf{A}^{(1)} ((\mathbf{A}^{(1)})^\dagger \mathbf{b})\|_2 / \sigma_{\mathbf{b}}$, resulting in 4,998 aggregated statistics of the asymptotic distribution, assuming convergence: mean, sd, skewness, kurtosis, min, max, as well as p1-p99, as presented in Figure 6 (depicting the intensity of its first two moments, mean and sd) and Table 2 (reporting a subset of all thirteen statistics for $m, p \bmod 10 = 0$).

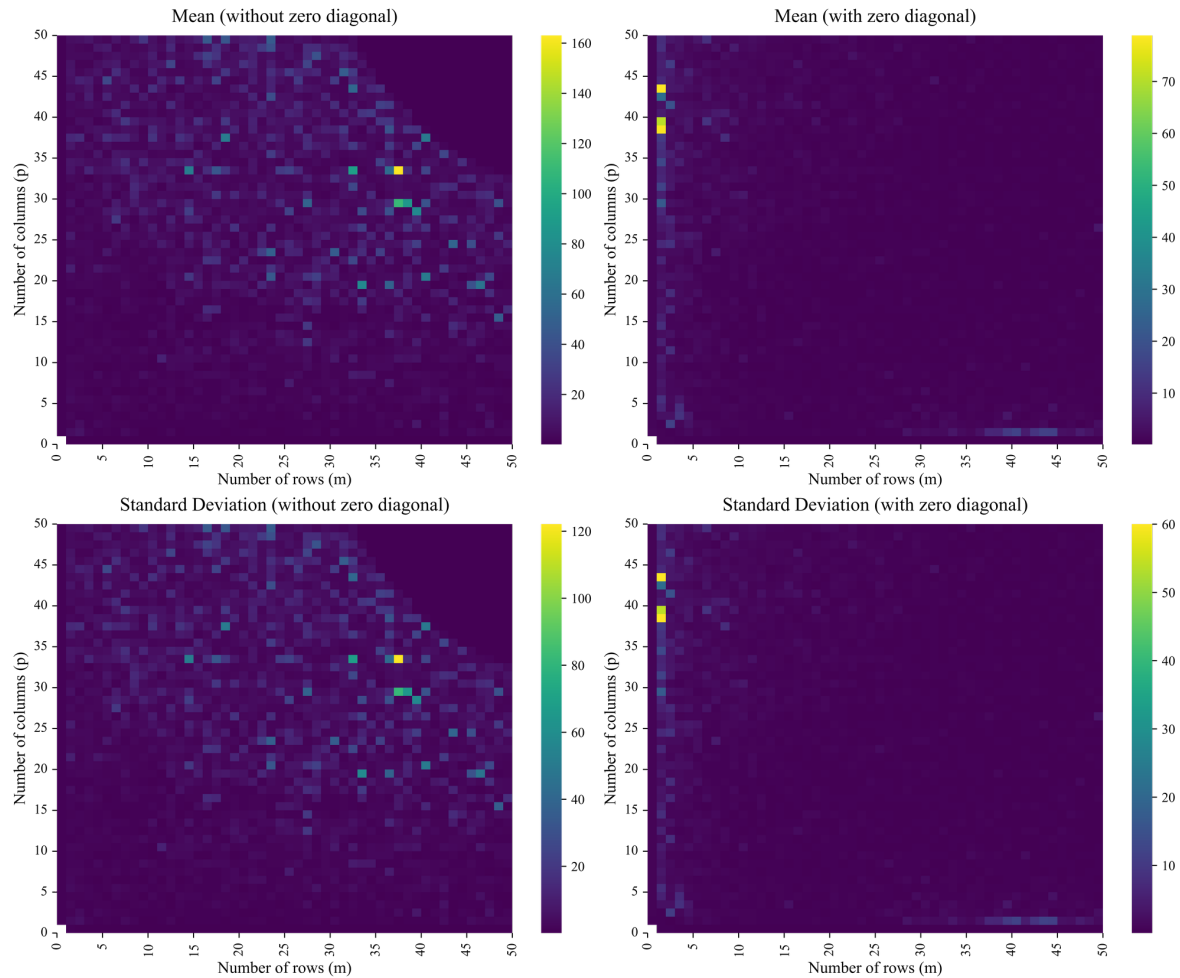


Figure 6. Heatmaps (based on `seaborn`'s `viridis` colormap) of means, $\text{mean} = \frac{1}{50,000} \sum_{i=1}^{50,000} \text{NRMSE}_i$, and standard deviations, $\text{sd} = \sqrt{\frac{1}{50,000-1} \sum_{i=1}^{50,000} (\text{NRMSE}_i - \text{mean})^2}$, computed from the NRMSE of $\hat{\mathbf{z}}^{(1)} = (\mathbf{A}^{(1)})^\dagger \mathbf{b}$, where $\mathbf{x} = \text{vec}(\mathbf{X} \in \mathbb{R}^{m \times p}, 1 \leq m, p \leq 50) \subseteq \mathbf{z}$, under two structural variants: with and without a zero diagonal ($\forall_{i \in \{1, \dots, \min(m, p)\}} \mathbf{X}_{ii} = 0$). Color reflects the magnitude of the statistic, with lighter shades indicating higher values. The figure was prepared using `matplotlib` in Python 3.

Table 2. Subset (at $m, p \bmod 10 = 0$) of the summary statistics for NRMSE from the Simulation Study.

m	p	mean	sd	skewness	kurtosis	min	max	p1	p5	p25	p50	p75	p95	p99
Without Zero Diagonal														
10	10	000.81	000.59	000.82	003.24	000.00	003.52	000.01	000.06	000.33	000.70	001.18	001.94	002.47
10	20	001.53	001.14	000.92	003.60	000.00	008.10	000.02	000.12	000.62	001.31	002.21	003.70	004.80
10	30	005.30	003.96	000.95	003.69	000.00	027.25	000.09	000.41	002.13	004.51	007.65	012.95	016.90
10	40	007.56	005.67	000.97	003.77	000.00	046.23	000.12	000.61	003.02	006.45	010.90	018.59	024.18
10	50	025.08	018.95	000.99	003.84	000.00	127.63	000.39	001.98	009.99	021.10	036.31	061.64	080.75
20	10	000.57	000.43	000.93	003.60	000.00	002.73	000.01	000.04	000.23	000.49	000.83	001.39	001.79
20	20	007.90	005.85	000.89	003.48	000.00	040.71	000.13	000.64	003.18	006.77	011.48	019.17	024.57
20	30	004.10	003.03	000.92	003.61	000.00	021.59	000.06	000.33	001.68	003.50	005.92	009.92	012.87
20	40	005.85	004.37	000.95	003.71	000.00	031.73	000.09	000.47	002.36	005.00	008.44	014.27	018.65
20	50	008.52	006.36	000.94	003.63	000.00	043.37	000.14	000.68	003.44	007.22	012.36	020.73	026.90
30	10	002.89	002.16	000.97	003.78	000.00	015.48	000.05	000.23	001.16	002.47	004.17	007.05	009.19
30	20	002.85	002.11	000.92	003.58	000.00	014.28	000.05	000.23	001.16	002.43	004.12	006.88	008.93
30	30	004.44	003.32	000.94	003.66	000.00	022.66	000.07	000.35	001.78	003.76	006.43	010.85	014.08
30	40	008.73	006.50	000.95	003.69	000.00	048.55	000.14	000.69	003.56	007.47	012.61	021.13	027.85
30	50	005.61	004.22	000.96	003.69	000.00	031.96	000.09	000.44	002.25	004.75	008.10	013.70	017.94
40	10	000.55	000.41	000.98	003.80	000.00	003.10	000.01	000.04	000.22	000.47	000.79	001.35	001.76
40	20	004.67	003.48	000.93	003.66	000.00	025.54	000.07	000.35	001.87	003.97	006.74	011.34	014.83
40	30	001.39	001.04	000.95	003.69	000.00	006.70	000.02	000.11	000.56	001.18	002.01	003.40	004.39
40	40	000.27	000.20	000.96	003.71	000.00	001.48	000.00	000.02	000.11	000.23	000.39	000.67	000.87
40	50	000.29	000.21	000.95	003.69	000.00	001.58	000.00	000.02	000.11	000.24	000.41	000.70	000.91
50	10	003.74	002.82	000.99	003.84	000.00	023.34	000.06	000.29	001.50	003.16	005.37	009.20	012.13
50	20	001.65	001.24	000.95	003.68	000.00	008.17	000.03	000.13	000.67	001.40	002.39	004.05	005.26
50	30	005.24	003.90	000.94	003.69	000.00	028.36	000.08	000.42	002.12	004.47	007.56	012.74	016.54
50	40	000.47	000.35	000.97	003.73	000.00	002.40	000.01	000.04	000.19	000.40	000.67	001.15	001.49
50	50	000.23	000.17	000.94	003.64	000.00	001.21	000.00	000.02	000.09	000.20	000.33	000.56	000.73
With Zero Diagonal														
10	10	000.71	000.52	000.85	003.37	000.00	003.20	000.01	000.06	000.29	000.62	001.04	001.71	002.20
10	20	000.69	000.51	000.92	003.60	000.00	003.54	000.01	000.05	000.28	000.59	001.00	001.67	002.16
10	30	001.93	001.44	000.95	003.74	000.00	010.17	000.03	000.15	000.78	001.64	002.77	004.69	006.13
10	40	004.65	003.50	000.97	003.76	000.00	023.68	000.07	000.37	001.86	003.92	006.72	011.43	014.89
10	50	001.94	001.47	000.98	003.84	000.00	012.42	000.03	000.15	000.77	001.64	002.81	004.75	006.25
20	10	000.70	000.52	000.92	003.59	000.00	003.70	000.01	000.06	000.28	000.60	001.02	001.70	002.20
20	20	000.70	000.52	000.93	003.57	000.00	003.68	000.01	000.05	000.28	000.60	001.01	001.71	002.21
20	30	000.42	000.31	000.94	003.68	000.00	002.24	000.01	000.03	000.17	000.36	000.61	001.02	001.34
20	40	000.27	000.20	000.94	003.64	000.00	001.27	000.00	000.02	000.11	000.23	000.39	000.65	000.85
20	50	002.06	001.55	000.97	003.81	000.00	011.72	000.03	000.16	000.83	001.75	002.98	005.04	006.64
30	10	000.45	000.34	000.95	003.67	000.00	002.40	000.01	000.04	000.18	000.38	000.65	001.10	001.43
30	20	000.42	000.32	000.95	003.71	000.00	002.16	000.01	000.03	000.17	000.36	000.61	001.03	001.34
30	30	000.45	000.34	000.96	003.76	000.00	002.53	000.01	000.03	000.18	000.38	000.65	001.10	001.43
30	40	000.34	000.26	000.97	003.73	000.00	001.88	000.01	000.03	000.14	000.29	000.50	000.85	001.10
30	50	000.57	000.43	000.96	003.75	000.00	003.21	000.01	000.05	000.23	000.48	000.82	001.39	001.81
40	10	000.34	000.26	000.96	003.76	000.00	001.83	000.01	000.03	000.14	000.29	000.49	000.84	001.08
40	20	000.33	000.25	000.98	003.77	000.00	001.68	000.01	000.03	000.13	000.28	000.47	000.80	001.05
40	30	000.63	000.47	000.97	003.80	000.00	003.36	000.01	000.05	000.25	000.53	000.90	001.52	001.99
40	40	000.38	000.28	000.98	003.82	000.00	002.23	000.01	000.03	000.15	000.32	000.54	000.92	001.20
40	50	000.30	000.23	000.97	003.78	000.00	001.76	000.00	000.02	000.12	000.26	000.44	000.75	000.98
50	10	000.61	000.46	000.97	003.74	000.00	003.16	000.01	000.05	000.24	000.52	000.88	001.48	001.95
50	20	000.17	000.13	000.97	003.76	000.00	000.87	000.00	000.01	000.07	000.14	000.24	000.41	000.54
50	30	000.32	000.24	000.97	003.82	000.00	002.10	000.00	000.03	000.13	000.28	000.47	000.80	001.03
50	40	000.61	000.46	000.95	003.67	000.00	003.04	000.01	000.05	000.24	000.52	000.89	001.49	001.95
50	50	000.68	000.51	000.96	003.70	000.00	003.76	000.01	000.05	000.27	000.58	000.98	001.66	002.16

Source: self-processed.

From the results of the Monte Carlo experiment, it is observable that: (1) the NRMSE from the first-step estimate $\hat{\mathbf{z}}^{(1)}$ and its distributional statistics exhibit increasing stability and boundedness as matrix dimensions $m \times p$ increase — specifically, for $m, p \gg 1$, both the mean and sd of NRMSE tend to converge, indicating improved estimator performance (e.g., under no diagonal restrictions, at $m = p = 5$, mean = 1.84 and sd = 1.29, while at $m = p = 45$, mean = 0.21 and sd = 0.16); (2) the zero-diagonal constraints ($\forall_{i \in \{1, \dots, \min(m, p)\}} \mathbf{X}_{ii} = 0$) reduce the degrees of freedom and lead to a much more uniform and dimensionally stable distribution of NRMSE across different matrix sizes (e.g., at $m = p = 5$, mean drops to 1.17 and sd to 0.85, while at $m = p = 45$, mean rises slightly to 0.26 and sd to 0.19); and (3) the distribution of NRMSE exhibits mild right skew and leptokurtosis — less so for the zero-diagonal case: under no diagonal restriction, the average skewness is 0.97 with a coefficient of

variation of 38.16% (i.e., > 0.00), and the average kurtosis is 4.15 with a variation of 304.61%; whereas under the zero-diagonal constraint, skewness averages 0.95 with only 8.57% variation, and kurtosis averages 3.76 with 25.94% variation (i.e., > 3.00). To sum up, the larger and less underdetermined (i.e., in an AP (TM) problem, the larger the $\mathbf{M}^{(1)}$ block of $\mathbf{A}^{(1)}$) the canonical system $\mathbf{A}^{(1)}\mathbf{z}^{(1)} = \mathbf{b}$, the lower the estimation error (i.e., the mean NRMSE) and its variance, and, provided a stable limiting law for NRMSE exists, the skewness and kurtosis from the simulations drift toward 0 and 3, respectively, consistent with — but not proving — a convergence toward a standard-normal distribution.

For an illustration of the applicability of the CLSP to the APs (TMs), CMLS (RPs), and LPRLS/QPRLS, the author-developed cross-platform Python 3 module, `pyclsp` (version $\geq 1.3.0$, available on PyPI) [98], — together with its interface co-modules for the APs (TMs), `pytmpinv` (version $\geq 1.2.0$, available on PyPI) [99], and the LPRLS/QPRLS, `pylppinv` (version $\geq 1.3.0$, available on PyPI) [100] (a co-module for the CMLS (RPs) is not provided, as the variability of $\forall_{i \in \{1, \dots, t\}} \mathbf{C}_i \equiv \mathbf{T}_i \mathbf{D}$ within the $\mathbf{C} = \begin{bmatrix} \mathbf{C}_1 \\ \vdots \\ \mathbf{C}_t \end{bmatrix}$, $u \geq t$, block of $\mathbf{A}^{(r)}$ prevents efficient generalization), — is employed for a sample (dataset) of i random elements $\mathcal{X}_i$ (scalars, vectors, or matrices depending on context) and their transformations, drawn independently from the standard normal distribution, i.e., $\{\mathcal{X}_1, \dots, \mathcal{X}_i\} \stackrel{\text{iid}}{\sim} \mathcal{N}(0, 1)$, under the mentioned random-variate seed 123456789 (configured for consistency with the above-described Monte Carlo experiment). Thus, using the cases of (a) a (non-negative symmetric) input-output table and (b) a (non-negative) zero-diagonal trade matrix as two **AP (TM) numerical examples**, this text simulates problems similar to the ones addressed in Pereira-López et al. [95] and Bolotov [16]: an underdetermined $\mathbf{X} \in \mathbb{R}^{m \times p}$ is estimated — from row sums, column sums, and k known values of two randomly generated matrices (a) $\mathbf{X}_{\text{true}} = 0.5 \cdot (|\mathcal{X}_1| + |\mathcal{X}_1|^\top)$ and (b) $\mathbf{X}_{\text{true},ij} = (|\mathcal{X}_1|)_{ij} \left(1 - \mathbb{1}_{\{i=j \leq \min(m,p)\}}\right)$, $(\mathcal{X}_1)_{ij} \stackrel{\text{iid}}{\sim} \mathcal{N}(0, 1)$, subject to (a) $\mathbf{X}_{ij} \geq 0$ and (b) $\mathbf{X}_{ij} \geq 0 \wedge \forall_{i \in \{1, \dots, \min(m,p)\}} \mathbf{X}_{ii} = 0$ — via CLSP assuming $\mathbf{x} = \text{vec}(\mathbf{X}) \subseteq \mathbf{z}$, $\mathbf{M} \equiv (\mathbf{I}_{mp})_{\mathcal{I}}$, $\mathbf{b} = \begin{bmatrix} \mathbf{b}_{\text{row sums}} \\ \mathbf{b}_{\text{column sums}} \\ \mathbf{b}_{\text{known values}} \end{bmatrix} = \begin{bmatrix} \mathbf{X}_{\text{true}} \mathbf{1} \\ (\mathbf{1}^\top \mathbf{X}_{\text{true}})^\top \\ (\text{vec}(\mathbf{X}_{\text{true}}))_{\mathcal{I}} \end{bmatrix}$, and $\mathcal{I} \sim \mathcal{U}\{1, \dots, mp\} \wedge |\mathcal{I}| = k$. The Python 3 code, based on `pytmpinv` and two other modules, installable by executing `pip install matplotlib numpy pytmpinv==1.2.0`, for (a) $m = p = 20$ and $k = 0.1 \cdot mp = 40$ and (b) $m = p = 40$ — estimated with the help of $m_{\text{subset}} \times n_{\text{subset}}$ -MNBLUE across $\left(\left\lceil \frac{m}{m_{\text{subset}}} \right\rceil \cdot \left\lceil \frac{n}{n_{\text{subset}}} \right\rceil\right) = \left(\left\lceil \frac{40}{20-1} \right\rceil \cdot \left\lceil \frac{40}{20-1} \right\rceil\right) = 9$ reduced models $m_{\text{reduced}} = p_{\text{reduced}} \leq 20$ assuming an estimability constraint of $mp \leq 400$ — and $k = 0.2 \cdot mp = 320$, with a non-iterated ($r = 1$), unique ($\alpha > 0$), and MNBLUE ($\alpha = 1$) (two-step) CLSP solution, is implemented in (a) Listing 1 and (b) Listing 2 (e.g., to be executed in a Jupyter Notebook).

Listing 1. Simulated numerical example for a symmetric input-output-table-based AP (TM) problem.

```
import numpy as np
from tmpinv import tmpinv
import matplotlib.pyplot as plt

seed = 123456789
rng = np.random.default_rng(seed)

# sample (dataset)
m, p = 20, 20 # matrix dimensions (m x p)
X_true = np.abs(rng.normal(size=(m, p))) # symmetric non-negative X_true
X_true = (X_true + X_true.T) / 2.0
idx = rng.choice(m * p, size=max(1, int(0.1 * (m * p))), # random numbers from
                  replace=False) # [0, (m * p) - 1], 10% of total

# model
M = np.eye(m * p)[idx, :] # unit matrix
b_row = X_true.sum(axis=1) # row sums (length m)
b_col = X_true.sum(axis=0) # column sums (length p)
b_val = X_true.reshape(-1, 1)[idx] # column vector
bounds = (0, None) # non-negativity
```

```

result = tmpinv(
    M=M, b_row=b_row, b_col=b_col, b_val=b_val,
    m=m, p=p, bounds=bounds, symmetric=True,
    r=1,                                     # a solution without refinement
    alpha=1.0                               # a unique MNBLUE estimator
)

# helpers
def fmt(v):
    v = float(v)
    return f"{v:.4e}" if abs(v) >= 1e6 or (v != 0 and abs(v) < 1e-4) else f"{v:.6f}"

# results
print("true_X:")
print(np.round(np.asarray(X_true), 4))
print("X_hat:")
print(np.round(result.x, 4))

print("\nDegrees of freedom:")
dof = result.model.A.shape[1] - np.linalg.matrix_rank(result.model.A)
p_null = round(dof / result.model.A.shape[1] * 100, 2)
print("nullity:", dof)
print("percentage of total unknowns:", str(p_null) + "%")

print("\nNumerical stability:")
print("kappaC:", fmt(result.model.kappaC))
print("kappaB:", fmt(result.model.kappaB))
print("kappaA:", fmt(result.model.kappaA))
corr = result.model.corr()
plt.figure(figsize=(8, 4))
plt.grid(True, linestyle="--", alpha=0.6)
plt.bar(range(len(corr["rmsa_i"])), corr["rmsa_i"])
plt.xlabel("Constraint index")
plt.ylabel("dRMSA (row deletion effect)")
plt.title(f"CLSP Correlogram (Total RMSA = {result.model.rmsa:.2f})")
plt.tight_layout()
plt.show()

print("\nGoodness-of-fit:")
x_true = X_true.flatten()
x_hat = result.x.flatten()
ss_res = np.sum((x_true - x_hat) ** 2)
ss_tot = np.sum((x_true - np.mean(x_true)) ** 2)
print("R2_user_defined:", fmt(1 - ss_res / ss_tot))
print("NRMSE:", fmt(result.model.nrmse))
print("Diagnostic band (min):", fmt(np.min(result.model.x_lower)))
print("Diagnostic band (max):", fmt(np.max(result.model.x_upper)))
print("Monte Carlo t-test:")
for kw, val in result.model.ttest(sample_size=30,
                                  seed=seed, distribution="normal", # NRMSE sample
                                  # seed and distribution
                                  ).items():
    print(f"{kw}: {float(val):.6f}")

```

In case (a) (see Listing 1 for code), the number of model (target) variables $\#\hat{\mathbf{X}}_M^*$ is $mp = 400$ with a nullity of $\dim(\ker(\mathbf{A}^{(1)})) = n - \text{rank}(\mathbf{A}^{(1)}) = 321$, corresponding to 80.25% of the total unknowns. Given the simulation of $\mathbf{X}_{\text{true}} \in \mathbb{R}^{20 \times 20}$, a matrix unknown in real-world applications, — i.e., CLSP is used to estimate the elements of an existing matrix from its row sums, $\mathbf{X}_{\text{true}}\mathbf{1} \in \mathbb{R}^{20}$, column sums, $(\mathbf{1}^\top \mathbf{X}_{\text{true}})^\top \in \mathbb{R}^{20}$, and a randomly selected 10% of its entries, $(\text{vec}(\mathbf{X}_{\text{true}}))_{\mathcal{I}} \in \mathbb{R}^{40}$, — the model's goodness of fit can be measured by a user-defined $R^2 \approx \max\left(0, 1 - \frac{\|\text{vec}(\mathbf{X}_{\text{true}}) - \hat{\mathbf{X}}_M^*\|_2^2}{\|\text{vec}(\mathbf{X}_{\text{true}}) - \frac{1}{400} \sum_{i=1}^{400} (\text{vec}(\mathbf{X}_{\text{true}})_i)\|_2^2}\right) \approx 0.278256$ (i.e., CLSP achieves an improvement of $\Delta R^2 = 0.178256$ over a hypothetical naïve predictor

reproducing the known 10% entries of $\mathbf{X}_{\text{true}, \mathcal{I}}$ and yielding a $R^2 \approx 0.1$, but still a modest value of R^2 , i.e. a relatively large error, $\text{NRMSE} = 12.649111$ with $\hat{\mathbf{x}}_M^*$ lying within wide condition-weighted diagnostic intervals $\hat{\mathbf{x}}_M^* \in [\min(\hat{\mathbf{x}}_{M, \text{lower}}^*), \max(\hat{\mathbf{x}}_{M, \text{upper}}^*)] = [-8.0847 \times 10^{15}, 8.0847 \times 10^{15}]$ reflecting the ill-conditioning of the strongly rectangular $\mathbf{A}^{(1)} \in \mathbb{R}^{80 \times 400}$, $\kappa(\mathbf{A}^{(1)}) = 4.9968 \times 10^{14}$, and only the left-sided Monte Carlo-based t -test for the mean of the NRMSE (on a sample of 30 NRMSE values obtained by substituting $\mathbf{b}$ with $(\mathbf{b}^{(\tau)})_i \stackrel{\text{iid}}{\sim} \mathcal{N}(0, 1)$) suggesting consistency with expectations (i.e., with the p-value ≈ 1). In terms of numerical stability, $\kappa(\mathbf{C}_{\text{canon}}) = 4.8023 \times 10^{14} \approx \kappa(\mathbf{A}^{(1)})$ with a low $\text{RMSA}_{i \in \{1, \dots, 40\}} = \text{RMSA} = 0.0400$, as confirmed by the correlogram produced in matplotlib, indicating a well-chosen constraints matrix, even given the underdeterminedness of the model (which also prevents a better fit).

Listing 2. Simulated numerical example for a (non-negative) trade-matrix-based AP (TM) problem.

```
import numpy as np
from tmpinv import tmpinv

seed = 123456789
rng = np.random.default_rng(seed)

# sample (dataset)
m, p = 40, 40
X_true = np.abs(rng.normal(size=(m, p)))
X_true = X_true * (1 - np.eye(m, p))
idx = rng.choice(m * p, size=max(1, int(0.2 * (m * p))),
                  replace=False)

# model
M = np.eye(m * p)[idx, :]
b_row = X_true.sum(axis=1)
b_col = X_true.sum(axis=0)
b_val = X_true.reshape(-1, 1)[idx]
bounds = (0, None)
result = tmpinv(
    M=M, b_row=b_row, b_col=b_col, b_val=b_val,
    m=m, p=p, zero_diagonal=True, reduced=(20, 20),
    bounds=bounds,
    r=1,
    alpha=1.0
)

# helpers
def fmt(v):
    v = float(v)
    return f"{v:.4e}" if abs(v) >= 1e6 or (v != 0 and abs(v) < 1e-4) else f"{v:.6f}"

# results
print("true_X:")
print(np.round(np.asarray(X_true), 4))
print("X_hat:")
print(np.round(result.x, 4))

print("\nDegrees of freedom:")
dof = np.array([CLSP.A.shape[1] - np.linalg.matrix_rank(CLSP.A)
                 for CLSP in result.model])
p_null = np.array([round(dof[i] / result.model[i].A.shape[1] * 100, 2)
                    for i in range(len(dof))])
print("nullity", np.min(dof).item(),
      np.max(dof).item())
print("percentage of total unknowns:", str(np.min(p_null).item()) + "%", "-",
      str(np.max(p_null).item()) + "%")
```

```

print("\nNumerical stability (min-max across models):")
kappaC = np.array([CLSP.kappaC for CLSP in result.model])
kappaB = np.array([CLSP.kappaB for CLSP in result.model])
kappaA = np.array([CLSP.kappaA for CLSP in result.model])
print("kappaC:", fmt(np.min(kappaC)), "-", fmt(np.max(kappaC)))
print("kappaB:", fmt(np.min(kappaB)), "-", fmt(np.max(kappaB)))
print("kappaA:", fmt(np.min(kappaA)), "-", fmt(np.max(kappaA)))

print("\nGoodness-of-fit (min-max across models):")
x_true = X_true.flatten()
x_hat = result.x.flatten()
mask = np.isfinite(x_true) & np.isfinite(x_hat)
if mask.any():
    ss_res = np.sum((x_true[mask] - x_hat[mask])**2)
    ss_tot = np.sum((x_true[mask] - np.mean(x_true[mask]))**2)
    nrmse = np.array([CLSP.nrmse for CLSP in result.model])
    x_lower = np.concatenate([np.array(CLSP.x_lower).reshape(-1) for CLSP in result.model])
    x_upper = np.concatenate([np.array(CLSP.x_upper).reshape(-1) for CLSP in result.model])
    print("R2_user_defined:", fmt(1 - ss_res / ss_tot))
    print("NRMSE:", fmt(np.min(nrmse)), "-", fmt(np.max(nrmse)))
    print("Diagnostic band (min):", fmt(np.min(x_lower)), "-", fmt(np.max(x_lower)))
    print("Diagnostic band (max):", fmt(np.min(x_upper)), "-", fmt(np.max(x_upper)))
    print("Monte Carlo t-test:")
    ttests = [CLSP.ttest(sample_size=30, seed=seed,
                        distribution="normal") for CLSP in result.model]
    keys = ttests[0].keys()
    for kw in keys:
        val = np.array([t[kw] for t in ttests], dtype=float)
        print(f"{kw}: {fmt(np.min(val))} - {fmt(np.max(val))}")

```

In case (b) (see Listing 2 for code), the corresponding number of model (target) variables in each of the $\left(\begin{bmatrix} 40 \\ 19 \end{bmatrix} \cdot \begin{bmatrix} 40 \\ 19 \end{bmatrix}\right) = 9$ reduced design submatrices $\mathbf{A}_{\text{subset}}^{(1)} \subset \mathbf{A}_{\text{reduced}}^{(1)}$, $\#\hat{\mathbf{x}}_{M,\text{reduced}}^*$, is $m_{\text{subset}}p_{\text{subset}} \leq 361$ — where $m_{\text{subset}} \leq 20 - 1 = 19$ and $p_{\text{subset}} \leq 20 - 1 = 19$, one row and one column being reserved for $\sum_{j \notin \mathcal{J}} a_{i,j}^{(1)}$ and $\sum_{i \notin \mathcal{I}} a_{i,j}^{(1)}$, which enter $\mathbf{A}_{\text{reduced}}^{(1)}$ as $\mathbf{S} = \begin{bmatrix} \mathbf{I}_{m-m_{\text{subset}}} & \mathbf{0} \\ \mathbf{0} & \mathbf{I}_{p-p_{\text{subset}}} \end{bmatrix}$ and the vector of slack variables $(\hat{\mathbf{y}}_{\text{reduced}}^*)^* = \begin{bmatrix} \hat{\mathbf{y}}_{\text{reduced}}^* \\ \hat{\mathbf{y}}_{\text{residual, reduced}}^* \end{bmatrix} \subset \hat{\mathbf{z}}_{\text{reduced}}^*$, so that $\mathbf{C}_{\text{reduced}}\mathbf{x}_{\text{reduced}} + \mathbf{S}_{\text{reduced}}\mathbf{y}_{\text{reduced}}^* = \mathbf{b}_{r,1:k}$ (i.e., the slack matrix compensates for the unaccounted row and column sums in the reduced models as opposed to the full one, such as case (a)) — with a nullity $\dim(\ker(\mathbf{A}_{\text{reduced}}^{(1)})) = n_{\text{reduced}} - \text{rank}(\mathbf{A}_{\text{reduced}}^{(1)}) \in [2, 290]$ (per reduced model), corresponding to 25.00%–72.68% of the total unknowns (per reduced model) — to compare, a full model, under the same inputs and no computational constraints (i.e., $mp \not\leq m_E p_E$), would have a nullity $\dim(\ker(\mathbf{A}^{(1)})) = n - \text{rank}(\mathbf{A}^{(1)}) = 1168$, corresponding to 73.00% of the total unknowns. In the examined example, — based on $\mathbf{X}_{\text{true}}\mathbf{1} \in \mathbb{R}^{40}$, $(\mathbf{1}^\top \mathbf{X}_{\text{true}})^\top \in \mathbb{R}^{40}$, and a randomly selected 20% of entries of the true matrix $\mathbf{X}_{\text{true}} \in \mathbb{R}^{40 \times 40}$, — the reduced-model block solution's goodness of fit could not be efficiently measured by a user-defined $R^2 \approx \max\left(0, 1 - \frac{\|\text{vec}(\mathbf{X}_{\text{true}}) - \text{vec}(\{\hat{\mathbf{x}}_{M,\text{reduced},1:m_{\text{subset}},1:n_{\text{subset}}}\})\|_2^2}{\|\text{vec}(\mathbf{X}_{\text{true}}) - \frac{1}{1600} \sum_{i=1}^{1600} (\text{vec}(\mathbf{X}_{\text{true}})_i)\|_2^2}\right) \approx 0$ (i.e., the block matrix constructed from $\left(\begin{bmatrix} 40 \\ 19 \end{bmatrix} \cdot \begin{bmatrix} 40 \\ 19 \end{bmatrix}\right) = 9$ reduced-model estimates led to $R^2 \approx \max(0, -5.435377)$, but to an error proportionate to the one in case (a), $\text{NRMSE}_{\text{reduced}} \in [3.464102, 15.748016]$ (per reduced model)) — in contrast, a full model would achieve $R^2 \approx \max\left(0, 1 - \frac{\|\text{vec}(\mathbf{X}_{\text{true}}) - \hat{\mathbf{x}}_M^*\|_2^2}{\|\text{vec}(\mathbf{X}_{\text{true}}) - \frac{1}{1600} \sum_{i=1}^{1600} (\text{vec}(\mathbf{X}_{\text{true}})_i)\|_2^2}\right) \approx 0.278427$, but at a cost of a greater error, $\text{NRMSE} = 29.427878$ — with $\hat{\mathbf{x}}_{M,\text{reduced}}^*$ lying within strongly varying condition-weighted diagnostic intervals $\hat{\mathbf{x}}_{M,\text{reduced}}^* \in [\min(\hat{\mathbf{x}}_{M,\text{lower, reduced}}^*), \max(\hat{\mathbf{x}}_{M,\text{upper, reduced}}^*)]$, where $\min(\hat{\mathbf{x}}_{M,\text{lower, reduced}}^*) \in [-330.638753, -2.5009 \times 10^{-19}]$ (per reduced model) and $\max(\hat{\mathbf{x}}_{M,\text{upper, reduced}}^*) \in [-2.5456 \times$

10^{-19} , 336.547225] (per reduced model), and varying results of Monte Carlo-based t -tests for the mean of the $\text{NRMSE}_{\text{reduced}}$ (on a sample of 30 $\text{NRMSE}_{\text{reduced}}$ values obtained by substituting $\mathbf{b}_r$ with $(\mathbf{b}_r^{(\tau)})_i \stackrel{\text{iid}}{\sim} \mathcal{N}(0, 1)$), where p -values range 4.8022×10^{-09} –1.000000 (left-sided), 0.000000–1.000000 (two-sided), and 3.7021×10^{-20} –1.000000 (right-sided) (per reduced model) — alternatively, a full model would lead to wider condition-weighted diagnostic intervals $\hat{\mathbf{x}}_M^* \in [\min(\hat{\mathbf{x}}_{M,\text{lower}}^*), \max(\hat{\mathbf{x}}_{M,\text{upper}}^*)] = [-1.5357 \times 10^{16}, 1.5357 \times 10^{16}]$ (i.e., reflecting the ill-conditioning of the strongly rectangular $\mathbf{A}^{(1)} \in \mathbb{R}^{400 \times 1600}$, $\kappa(\mathbf{A}^{(1)}) = 2.2196 \times 10^{14}$) and only the left-sided Monte Carlo-based t -test for the mean of the NRMSE (on a sample of 30 NRMSE values obtained by substituting $\mathbf{b}$ with $(\mathbf{b}^{(\tau)})_i \stackrel{\text{iid}}{\sim} \mathcal{N}(0, 1)$) suggesting consistency with expectations (i.e., with the p -value ≈ 1). In terms of numerical stability, $\kappa(\mathbf{C}_{\text{canon, reduced}}) \in [2.236068, 6.244998] < \kappa(\mathbf{A}_{\text{reduced}}^{(1)}) \in [4.509742, 9.325735]$ (per reduced model), which indicates well-conditioning of all the reduced models — conversely, in a full model, $\kappa(\mathbf{C}_{\text{canon}}) = 1.0907 \times 10^{14} < \kappa(\mathbf{A}^{(1)})$ (therefore, a full model ensures an overall better fit but a lower fit quality, i.e., a trade-off).

As the **CMLS (RP) numerical example**, this text addresses a problem similar to the one solved in Bolotov [15,96]: a coefficient vector $\boldsymbol{\beta}$ from a (time-series) linear regression model in its classical (statistical) notation $y_t = \mathbf{x}_t \boldsymbol{\beta} + \epsilon_t$, s.t. $\sum_{j=1}^{i+1} \beta_j = c \wedge \forall_{t \in \mathbf{P}} \Delta y_t = 0 \wedge \forall_{t \in \mathbf{P}} \Delta^2 y_t \geq 0$ with $\Delta^n y_t$ denoting n -th order differences (i.e., the discrete analogue of $\frac{d^n y_t}{dt^n}$) — where y_t is the dependent variable, $\mathbf{x}_t = [1, x_{t,1}, \dots, x_{t,i}]$ is the vector of regressors with a constant, ϵ_t is the model's error (residual), c is a constant, and $\mathbf{P} \subseteq \{t_j : j = 1, \dots, k\}$ is the set of stationary points — is estimated on a simulated sample (dataset) with the help of CLSP assuming $\mathbf{x}_M \equiv \boldsymbol{\beta} \in \mathbb{R}^p$ and $\mathbf{D} \equiv \mathbf{X} = [\mathbf{1}, \mathbf{x}_1, \dots, \mathbf{x}_i] \in \mathbb{R}^{k \times p}$, $p = i + 1$. The Python 3 code, based on numpy and pylsp modules, installable by executing `pip install numpy pylsp==1.3.0`, for $k = 500$, $p = 6$, $c = 1$, $\mathbf{y} = \mathbf{D} \boldsymbol{\beta}_{\text{true}} + \boldsymbol{\epsilon}$, where $\forall_{j \in \{2, \dots, 6\}} \mathbf{D}_{\cdot, j} \equiv \mathcal{X}_j$, $(\mathcal{X}_j)_i \stackrel{\text{iid}}{\sim} \mathcal{N}(0, 1)$, $\mathcal{X}_j \in \mathbb{R}^{500}$, $\boldsymbol{\beta}_{\text{true}} = \frac{\mathcal{X}_\beta}{1^\top \mathcal{X}_\beta}$, $(\mathcal{X}_\beta)_i \stackrel{\text{iid}}{\sim} \mathcal{N}(0, 1)$, $\mathcal{X}_\beta \in \mathbb{R}^6$, $\boldsymbol{\epsilon} \equiv \mathcal{X}_\epsilon$, $(\mathcal{X}_\epsilon)_i \stackrel{\text{iid}}{\sim} \mathcal{N}(0, 1)$, $\mathcal{X}_\epsilon \in \mathbb{R}^{500}$, and $\mathbf{P} = \{t : \Delta \mathbf{y} = 0\}$, with a non-iterated ($r = 1$), unique ($\alpha > 0$), and MNBLUE ($\alpha = 1$) two-step CLSP solution (consistent with the ones from cases (a) and (b) for APs (TMs)), is implemented in Listing 3 (e.g., for a Jupyter Notebook).

Listing 3. Simulated numerical example for a (time-series) stationary-points-based CMLS (RP) problem.

```
import numpy as np
from clsp import CLSP

seed = 123456789
rng = np.random.default_rng(seed)

# sample (dataset)
k = 500 # number of observations in D
p = 6 # number of regressors
c = 1 # sum of coefficients
D = np.empty((k, p))
D[:, 0] = 1.0 # constant
D[:, 1:] = rng.normal(size=(k, p - 1)) # D_{\cdot, j} ~ N(0, 1), 2 <= j <= p
b_true = rng.normal(size=p) # b_true ~ N(0, 1)
b_true = (b_true / b_true.sum()) * c
e = rng.normal(size=(k, 1)) # e ~ N(0, 1)
y = (D @ b_true).reshape(-1, 1) + e # y_t = D @ b_true + e_t

# model
b = np.vstack([
    np.asarray([c]), # c (the sum of coefficients)
    np.zeros((k - 2, 1)), # zeros
    np.zeros((k - 1, 1)), # zeros
    y # values of y_t
])
```

```

C      = np.vstack([
            np.ones((1, p)),                # a row of ones
            np.diff(D, n=2, axis=0),         # the 2nd differences
            np.diff(D, n=1, axis=0)          # the 1st differences
        ])
S      = np.block([
            [np.zeros((1, k-2))],            # a zero vector
            [np.diag(np.sign(np.diff(y.ravel(), n=2)))], # a diagonal sign matrix
            [np.zeros((k-1, k-2))],          # a zero matrix
        ])
model  = CLSP().solve(
            problem="cmls", b=b, C=C, S=S, M=D,
            r=1,                                # a solution without refinement
            alpha=1.0                            # a unique MNBLUE estimator
        )

# results
print("true_beta_(x_M):")
print(np.round(np.asarray(b_true).flatten(), 4))

print("beta_hat_(x_M_hat):")
print(np.round(model.x.flatten(), 4))

print("\nNumerical_stability:")
print("kappaC:", round(model.kappaC, 4))
print("kappaB:", round(model.kappaB, 4))
print("kappaA:", round(model.kappaA, 4))

print("\nGoodness-of-fit:")
print("R2_partial:", round(model.r2_partial, 6))
print("NRMSE_partial:", round(model.nrmse_partial, 6))
print("Diagnostic_band_(min):", np.round(np.min(model.x_lower), 4))
print("Diagnostic_band_(max):", np.round(np.max(model.x_upper), 4))
print("MonteCarlo_t-test:")
for kw, val in model.ttest(sample_size=30,                # NRMSE_partial sample
                           seed=seed, distribution="normal", # seed and distribution
                           partial=True).items():
    print(f"{kw}: {float(val):.6f}")

```

Compared to the true values $\beta_{\text{true}} = [0.1687, -0.3593, 0.5969, -0.0628, -0.1742, 0.8307]$, the CMLS (RP) estimate is $\hat{\beta} = \hat{\mathbf{x}}_M^* = [0.2024, -0.1121, 0.2095, -0.0396, -0.0296, 0.2611]$ with a modest $R^2_{\text{partial}} = 0.293328$ (i.e., a greater error, $\text{NRMSE}_{\text{partial}} = 0.840638$), moderate condition-weighted diagnostic intervals $\hat{\mathbf{x}}_M^* \in [-29.7780, 30.3003]$, and only the right-sided Monte Carlo-based t -test for the mean of the $\text{NRMSE}_{\text{partial}}$ (on a sample of 30 $\text{NRMSE}_{\text{partial}}$ values obtained by substituting $\mathbf{b}$ with $(\mathbf{b}^{(\tau)})_i \stackrel{\text{iid}}{\sim} \mathcal{N}(0, 1)$) — in the example under scrutiny, $\text{NRMSE}_{\text{partial}}$ is preferable to NRMSE due to $\hat{\mathbf{y}}^* \gg \hat{\mathbf{x}}^*$ — suggesting consistency with expectations (i.e., with the p -value ≈ 1). Similarly, in terms of numerical stability, $\kappa(\mathbf{C}_{\text{canon}}) = 138.6109 > \kappa(\mathbf{A}^{(1)}) = 127.8287$, indicating that the constraint block is ill-conditioned, most likely, because of the imposed data- (rather than theory-) based definition of stationary points, which rendered Step 2 infeasible ($\hat{\mathbf{x}}_M^* = \hat{\mathbf{x}}_M^{(1)}$) (limiting the fit quality).

Finally, in the case of a **LPRLS/QPRLS numerical example**, this text simulates potentially ill-posed LP (QP) problems, similar to the ones addressed in Whiteside et al. [93] and Blair [76,94]: a solution $\hat{\mathbf{x}}^* \geq 0$ in its classical LP (QP) notation $\mathbf{x}^* = \arg \max_{\mathbf{x} \in \mathbb{R}^p} \mathbf{c}^\top \mathbf{x}$ ($\mathbf{x}^* = \arg \min_{\mathbf{x} \in \mathbb{R}^p} \left\{ \frac{1}{2} \mathbf{x}^\top \mathbf{Q} \mathbf{x} + \mathbf{c}^\top \mathbf{x} \right\}$), s.t. $\mathbf{A} \mathbf{x} \leq \mathbf{b} \equiv \mathbf{A}_{ub} \mathbf{x} \leq \mathbf{b}_{ub} \wedge \mathbf{A}_{eq} \mathbf{x} = \mathbf{b}_{eq}$ — where $\mathbf{Q} \in \mathbb{R}^{p \times p}$ is a symmetric positive definite matrix, $\mathbf{c} \in \mathbb{R}^p$, $\mathbf{A} = \begin{bmatrix} \mathbf{A}_{ub} \\ \mathbf{A}_{eq} \end{bmatrix} \in \mathbb{R}^{(m_{ub}+m_{eq}) \times p}$, and $\mathbf{b} = \begin{bmatrix} \mathbf{b}_{ub} \\ \mathbf{b}_{eq} \end{bmatrix} \in \mathbb{R}^{m_{ub}+m_{eq}}$ — is estimated from two randomly generated (coefficient) matrices $\mathbf{A}_{ub} \equiv \mathcal{X}_1, (\mathcal{X}_1)_{i,j} \stackrel{\text{iid}}{\sim} \mathcal{N}(0, 1)$, and $\mathbf{A}_{eq} \equiv \mathcal{X}_2, (\mathcal{X}_2)_{i,j} \stackrel{\text{iid}}{\sim} \mathcal{N}(0, 1)$, and two (right-hand-side) vectors $\mathbf{b}_{ub} \equiv \mathcal{X}_3, (\mathcal{X}_3)_i \stackrel{\text{iid}}{\sim} \mathcal{N}(0, 1)$, and $\mathbf{b}_{eq} \equiv \mathcal{X}_4, (\mathcal{X}_4)_i \stackrel{\text{iid}}{\sim} \mathcal{N}(0, 1)$, via CLSP

assuming $\mathbf{C} = \begin{bmatrix} \mathbf{A}_{ub} \\ \mathbf{A}_{eq} \end{bmatrix}$, $\mathbf{S} = \mathbf{Q}_S^L \begin{bmatrix} \mathbf{S}_s \subseteq \mathbf{I}_{m_{ub}} \\ \mathbf{0} \end{bmatrix} \mathbf{Q}_S^R$, and $\mathbf{b} = \begin{bmatrix} \mathbf{b}_{ub} \\ \mathbf{b}_{eq} \end{bmatrix}$, where $\mathbf{Q}_S^L$ and $\mathbf{Q}_S^R$ are permutation matrices, while omitting $\mathbf{Q}$ and $\mathbf{c}$. The Python 3 code, based on numpy and pylppinv modules, installable by executing `pip install numpy pylppinv==1.3.0`, for $m_{ub} = 50$, $m_{eq} = 25$, and $p = 500$, with a non-iterated ($r = 1$), unique ($\alpha > 0$), and MNBLUE ($\alpha = 1$) (two-step) CLSP solution (analogously consistent with the ones from cases (a) and (b) for APs (TMs) and the one from CMLS (RPs)), is implemented in Listing 4 (e.g., for a Jupyter Notebook).

Listing 4. Simulated numerical example for a underdetermined and potentially infeasible LP (QP) problem.

```

import numpy as np
from lppinv import lppinv

seed = 123456789
rng = np.random.default_rng(seed)

# sample (dataset)
A_ub = rng.normal(size=(50, 500)) # underdetermined LP/QP matrix
A_eq = rng.normal(size=(25, 500)) # underdetermined LP/QP matrix
b_ub = rng.normal(size=(50, 1)) # may be inconsistent with A_ub
b_eq = rng.normal(size=(25, 1)) # may be inconsistent with A_eq
model = lppinv(
    A_ub=A_ub, A_eq=A_eq, b_ub=b_ub, b_eq=b_eq,
    non_negative=False, # allow negative values
    r=1, # a solution without refinement
    alpha=1.0 # a unique MNBLE estimator
)

# results
print("x_hat(x_M_hat):")
print(np.round(model.x.flatten(), 4))

print("\nDegrees of freedom:")
dof = model.A.shape[1] - np.linalg.matrix_rank(model.A)
p_null = round(dof / model.A.shape[1] * 100, 2)
print("nullity:", dof)
print("percentage of total unknowns:", str(p_null) + "%")

print("\nNumerical stability:")
print("kappaC:", round(model.kappaC, 4))
print("kappaB:", round(model.kappaB, 4))
print("kappaA:", round(model.kappaA, 4))

print("\nGoodness-of-fit:")
print("NRMSE:", round(model.nrmse, 6))
print("Diagnostic band(min):", np.round(np.min(model.x_lower), 4))
print("Diagnostic band(max):", np.round(np.max(model.x_upper), 4))
print("Monte Carlo t-test:")
for kw, val in model.ttest(sample_size=30, # NRMSE sample
                           seed=seed, distribution="normal", # seed and distribution
                           ).items():
    print(f"{kw}: {float(val):.6f}")

```

The nullity (i.e., underdeterminedness) of the CLSP design matrix, $\mathbf{A}^{(1)}$, is $\dim(\ker(\mathbf{A}^{(1)})) = n - \text{rank}(\mathbf{A}^{(1)}) = 475$, corresponding to 86.36% of the total unknowns, accompanied by a greater (relatively high) error, $\text{NRMSE} = 12.247449$, moderate condition-weighted diagnostic intervals $\hat{\mathbf{x}}_M^* \in [-1.2919, 1.3918]$, and only the right-sided Monte Carlo-based t -test for the mean of the NRMSE (on a sample of 30 NRMSE values obtained by substituting $\mathbf{b}$ with $(\mathbf{b}^{(\tau)})_i \stackrel{\text{iid}}{\sim} \mathcal{N}(0, 1)$) suggesting consistency with expectations (i.e., with the p-value ≈ 1). Similarly, in terms of numerical stability, $\kappa(\mathbf{C}_{\text{canon}}) = 2.2168 = \kappa(\mathbf{A}^{(1)}) = 2.2168$, indicating that the constraint block is well-conditioned despite

a strongly rectangular $\mathbf{A}^{(1)} \in \mathbb{R}^{75 \times 550}$. Hence, CLSP can be efficiently applied to AP (TM), CMLS (RP), and LPRLS/QPRLS special cases using Python 3, and the reader may wish to experiment with the code in Listings 1–4 by relaxing its uniqueness (setting $\alpha = 0$) or the MNBLUE characteristic (setting $\alpha \in [0, 1)$).

8. Discussion and Conclusions

The proposed modular two-step Convex Least Squares Programming (CLSP) estimator attempts to supersede historical efforts developed mainly in the 1960s–1990s to address ill-posed problems in both LS and LP (QP) — namely, the mentioned Wolfe [84], Stoer [80], Waterman [81], Dax [85], Osborne [78], George and Osborne [79], Barnes et al. [86], Donoho and Tanner [77], Grafarend and Awange [82], Bei et al. [87], and Qian et al. [83] — by correcting (1) the estimation, based on a Moore-Penrose pseudoinverse $\mathbf{A}^+$ (which has already been applied in an ad hoc manner in the field, particularly in econometrics in the 2010s; consult Pereira-López et al. [95] and Bolotov [16]) and a Bott-Duffin inverse $\mathbf{A}_Z^+$ (based on the literature research, this text may be the first to consider a constrained generalized inverse to generalize an $\mathbf{A}^+$ -based estimator) with the help of a (2) regularization-based correction [4] (pp. 477–479), achieving greater fit quality in terms of ensuring strict satisfaction of problem constraints (previously infeasible under an $\mathbf{A}^+$ -based solution), while yielding a (potentially) unique and MNBLUE solution (by setting $\alpha = 1$). At the same time, CLSP is well grounded in the algorithmic logic of the works in question, such as Wolfe [84] and Dax [85], as further formalized by, e.g., Osborne [78]. However, the primary strength of the framework in question lies in its “generality”: in contrast to the (individual) problem-specific algorithms, such as the ones developed by Übi [88,89,90,91,92] — routines combining (NN)LS, LP, and QP within a single framework to solve both well- and ill-posed convex optimization problems, — CLSP provides a unified estimator with a proper methodology for numerical stability and goodness-of-fit assessment, centered around RMSA and NRMSE. A secondary strength is its practical “versatility”: through at least three classes of special cases — APs (TMs), CMLS (RPs), and LPRLS/QPRLS, — it allows estimation of (rectangular) input-output tables (national, inter-country, or satellite), structural matrices (such as trade, country-product, investment, or migration data), and previously infeasible textbook-definition models of economic variables and/or LP (QP) problems. Furthermore, the inclusion of $\mathbf{A}_Z^+$ allows probabilistic and Bayesian applications, since its subspace restrictions can be interpreted in terms of random variables, prior distributions, as well as Monte Carlo-based inference. To conclude, CLSP has already been fully implemented in Python 3 (pyclsp, pytmpinv, and pylppinv) and is to be ported into R and Stata in the future, making it a readily available alternative to purely theoretical or not freely accessible algorithms or proprietary packages, and applicable in a real-world setting (such as econometric research, replacing [15,16,95] with the two-step corrected estimation).

Author Contributions: The manuscript was prepared by a sole author.

Funding: This research received no external funding.

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Data Availability Statement: The numerical examples used in the manuscript are based on simulated data. The Python 3.10 and Stata 14 code for replication purposes is provided in Listings 1–4 and B.1.

Acknowledgments: All arguments, structure, and conclusions in this manuscript are the author’s own. During the preparation of the manuscript, the author resorted to ChatGPT-5 (OpenAI) solely to improve clarity and language, including refinement of style, formal precision in proofs of theorems, and assistance with interpretation of data and results. No content was adopted without the author’s review and editing, and the author claims full responsibility for the final content of the manuscript.

Conflicts of Interest: The author declares no conflicts of interest.

Abbreviations

The following abbreviations are used in this manuscript (in order of their appearance in the sections):

CLSP	Convex Least Squares Programming
LS	Least Squares
OLS	Ordinary Least Squares
NNLS	Non-Negative Least Squares
SVD	Singular Value Decomposition
LP	Linear Programming
QP	Quadratic Programming
Lasso	Least Absolute Shrinkage and Selection Operator
Ridge	Ridge Regression (Tikhonov regularization)
MNBLUE	Minimum-Norm Best Linear Unbiased Estimator
BLUE	Best Linear Unbiased Estimator
RMSA	Root Mean Square Alignment
RMSE	Root Mean Square Error
NRMSE	Normalized Root Mean Square Error
ANOVA	Analysis of Variance
CLT	Lindeberg-Lévy Central Limit Theorem
APs	Allocation Problems
TMs	Tabular Matrix Problems
UNCTAD	UN Trade and Development
CMLS	Constrained-Model Least Squares
RP	Regression Problems
NBER	U.S. National Bureau of Economic Research
LPRLS	Linear Programming via Regularized Least Squares
QPRLS	Quadratic Programming via Regularized Least Squares
iid	independent and identically distributed

Appendix A

This appendix presents (a) an equation of the Bott-Duffin inverse, $(\mathbf{A}_Z^{(r)})^\dagger$, based on the singular value decomposition (SVD), as well as (b) a partitioned decomposition of $\mathbf{B}^{(r)}$.

Appendix A.1

Let $\mathbf{A}^{(r)} \in \mathbb{R}^{m \times n}$ be a real matrix, and let $\mathbf{Z} \in \mathbb{R}^{n \times n}$ be a symmetric idempotent matrix (i.e., $\mathbf{Z}^2 = \mathbf{Z} = \mathbf{Z}^\top$) defining a subspace restriction. Then, the Bott-Duffin inverse $(\mathbf{A}_Z^{(r)})^\dagger \in \mathbb{R}^{n \times m}$ is given by the formula $(\mathbf{A}_Z^{(r)})^\dagger = \left(\mathbf{Z}(\mathbf{A}^{(r)})^\top \mathbf{A}^{(r)} \mathbf{Z} \right)^\dagger \mathbf{Z}(\mathbf{A}^{(r)})^\top$, as defined in Equation (5), where the projected Gram matrix $\mathbf{Z}(\mathbf{A}^{(r)})^\top \mathbf{A}^{(r)} \mathbf{Z}$ is computable via its SVD:

$$\mathbf{Z}(\mathbf{A}^{(r)})^\top \mathbf{A}^{(r)} \mathbf{Z} = \mathbf{U} \mathbf{\Sigma} \mathbf{V}^\top = \mathbf{U} \mathbf{\Sigma} \mathbf{U}^\top \quad (\text{A1})$$

where $\mathbf{U}, \mathbf{V} \in \mathbb{R}^{n \times n}$ are orthogonal (since the projected Gram matrix is symmetric positive semidefinite, $\mathbf{U} = \mathbf{V}$), and $\mathbf{\Sigma} = \text{diag}(\sigma_1, \dots, \sigma_{\text{rank}}, 0, \dots, 0) \in \mathbb{R}^{n \times n}$ contains the singular values. Hence, for $\mathbf{\Sigma}^\dagger = \text{diag}(\sigma_1^{-1}, \dots, \sigma_{\text{rank}}^{-1}, 0, \dots, 0)$, $\mathbf{U} \mathbf{\Sigma}^\dagger \mathbf{U}^\top$ substitutes $\left(\mathbf{Z}(\mathbf{A}^{(r)})^\top \mathbf{A}^{(r)} \mathbf{Z} \right)^\dagger$:

$$(\mathbf{A}_Z^{(r)})^\dagger = \left(\mathbf{Z}(\mathbf{A}^{(r)})^\top \mathbf{A}^{(r)} \mathbf{Z} \right)^\dagger \mathbf{Z}(\mathbf{A}^{(r)})^\top = \left(\mathbf{U} \mathbf{\Sigma}^\dagger \mathbf{U}^\top \right) \mathbf{Z}(\mathbf{A}^{(r)})^\top \quad (\text{A2})$$

Appendix A.2

Let $\text{rank}([\mathbf{C} \ \mathbf{S}]) \leq \text{rank}(\mathbf{C}) + \text{rank}(\mathbf{S})$ in the analysis of numerical stability of the first estimate $\hat{\mathbf{z}}^{(r)}$, as defined in Equation (2). Then a decomposition of $\mathbf{B}^{(r)}$ reveals how the interaction of constraints $\mathbf{C}$ and slack structure $\mathbf{S}$ forms the curvature of the solution space:

$$\begin{aligned} \mathbf{B}^{(r)} &= \mathbf{A}^{(r)} \begin{bmatrix} \mathbf{C} & \mathbf{S} \end{bmatrix}^{\dagger} \\ &= f\left(\begin{bmatrix} \mathbf{C} & \mathbf{S} \end{bmatrix}, g\left(\mathbf{M}, \mathbf{b} - \mathbf{A}^{(r-1)} \hat{\mathbf{z}}^{(r-1)}\right)\right) \cdot \begin{bmatrix} \mathbf{C}^{\top} \mathbf{C} & \mathbf{C}^{\top} \mathbf{S} \\ \mathbf{S}^{\top} \mathbf{C} & \mathbf{S}^{\top} \mathbf{S} \end{bmatrix}^{\dagger} \begin{bmatrix} \mathbf{C}^{\top} \\ \mathbf{S}^{\top} \end{bmatrix} \end{aligned} \tag{A3}$$

While the block $\mathbf{C}^{\top} \mathbf{C}$ captures the contribution of the constraints and $\mathbf{S}^{\top} \mathbf{S}$ reflects the curvature introduced by slack variables, the off-diagonal blocks $\mathbf{C}^{\top} \mathbf{S}$ and $\mathbf{S}^{\top} \mathbf{C}$ encode cross-interactions between constraints and slacks: if these off-diagonal blocks vanish, the effects are orthogonal and additive (in the column space sense since $\mathbf{C}^{\top} \mathbf{S} = 0$ implies $\mathcal{R}(\mathbf{C}) \perp \mathcal{R}(\mathbf{S})$); otherwise, $\mathbf{C}$ is orthogonalized with respect to $\mathcal{R}(\mathbf{S})$, which is effected by $(\mathbf{I} - \mathbf{S} \mathbf{S}^{\dagger}) \mathbf{C}$. Hence, the presence of $\mathbf{S}$ in $[\mathbf{C} \ \mathbf{S}]$ can improve or worsen the ill-posedness of the canonized problem: namely, $\mathbf{C}^{\top} \mathbf{S} \neq 0$ tends to worsen the condition number $\kappa([\mathbf{C} \ \mathbf{S}])$.

Appendix B

This appendix presents Stata code, i.e., a .do-file, for the Monte Carlo experiment, compatible with version 14.0 and above (Stata/MP with at least four CPUs is recommended for replication — specifically, 50,000 iterations require substantial CPU time and memory).

Stata/MP 14.0 dependencies:

TITLE

’SUMMARIZEBY’: module to use statsby functionality with summarize

DESCRIPTION/AUTHOR(S)

This routine combines the `summarize` command with `statsby`. It iteratively collects `summarize` results for each variable with `statsby` and saves the result to memory or to a Stata .dta file.

KW: data management

KW: summarize

KW: statsby

Requires: Stata version 8

Distribution–Date: 20221012

Author: Ilya Bolotov, Prague University of Economics and Business

Support: email `ilya.bolotov@vse.cz`

INSTALLATION FILES

`summarizeby.ado`

`summarizeby.sthlp`

(type `ssc install summarizeby` to install)

Appendix B.1

Listing B.1. Stata code from the Monte Carlo experiment dedicated to mapping the distribution of NRMSE.

```
*! version 1.0.2 07oct2022
```

```

version 14.0

***# Bookmark #1
* ===== CHECK FOR MISSING SOFTWARE ===== *

***# Bookmark #2
* ===== CONFIGURATION ===== *
clear all
tempfile tmpf DTA_FILE // temporary file
tempname t // temporary frame
glo date : di %tdCYND daily("$S_DATE", "DMY") // date
glo date = cond("`1'" != "", "`1'", "$date")
glo compiled_by "Ilya Bolotov, MBA, Ph.D."
// load a pre-save (if applicable)
cap use "AP____DATA_pre-save.dta", clear
// create folder and/or open log
cap log using "AP____DATA_$date.smcl", name(data) replace

***# Bookmark #3
* ===== MATA CODE ===== *
sca `t' = clock(c(current_date) + " " + c(current_time), "DMYhms")
// Experiment, QUAD PRECISION
mata: d = (1::50) // dimensions of 'a'
mata: n = 50000 // iterations

loc RS real scalar
loc RV real colvector
loc RM real matrix
loc SS string scalar
loc TM transmorphic matrix
loc PF pointer('RM' function) scalar
loc VV void

loc version "version 14:"

mata:
mata set matastrict on

`RS' appinv_dist(`RS' r, `RS' c, `RS' n, |`RS' f_d, `RS' f_n, `RS' tol,
                `RS' seed, `PF' dist) // solver */
{
    `RV' e, i
    `RM' a, b
    `TM' tmp, f

    // general configuration
    f_d = f_d != . ? f_d : 0 // flag: zero-diagonal */
    f_n = f_n != . ? f_n : 0 // flag: uniform|normal */
    tol = tol != . ? tol : c("epsdouble") // tolerance */
    seed = seed != . ? seed : 123456789 // MC seed, distribution */
    dist = dist ? dist : (f_n ? &appinv_rnormal() : &appinv_runiform())

    // build the dataset (if the observation doesn't exist)
    if (st_nvar()) ///
    if (sum(rowsum(st_data(., ("r","c","zerod")) :== (r,c,f_d)) :== 3)) ///
    return(0)
    else (r,c,f_d)
    // TM problem, 'a' and 'b'
    a = I(r)#J(1,c,1)\J(1,r,1)#I(c)\
        (f_d ? (tmp=(I(min((r, c)))#(1,J(1,c,0)))[, (1..min((r,c))*c)],
            J(rows(tmp),r*c-cols(tmp),0)
            : J(0,r*c,.)) // diagonal of 'a' -> 0 */
    /* format: %e[2].Ofc */
    f = (tmp=trunc(floatround(log10(n)))) + trunc(floatround(tmp / 3)) + 2

```

```

printf("\n{txt}Simulations ({res}%"+ strofreal(f - 1)+".0fc{txt})\n", n)
printf("----- 1 ----- 2 ----- 3 ----- 4 ----- 5\n")
/* perform MC with n simulations */
rseed(seed)
for (i=1;i<(n+1);i++) {
    e = e\sqrt(quadcross((tmp=(b>(*dist)(r+c + (f_d ? min((r,c)) : 0), 1)) -
        a * svsolve(a, b, ., tol)), tmp) / r / variance(b))
    if (! mod(i, 5)) { printf("....."); displayflush(); }
    if (! mod(i, 50)) printf("%"+strofreal(f)+".0fc\n", i)
}
/**/
(void) stata("preserve", 1)
(void) st_dropvar(.) /* e -> summarizeby */
(void) st_addvar("double", "NRMSE", 1)
(void) st_addobs(rows(e), 1)
(void) st_store(., ("NRMSE"), e)
(void) stata("summarizeby, clear d", 1)
(void) st_addvar("long", ("r","c","zerod","rank","nvar"))
(void) st_store(st_nobs(), ("r","c","zerod","rank","nvar"),
    (r,c,f_d,rank(a),(f_d ? r*c - min((r,c)) : r*c)))
(void) stata("save 'tmpf', replace", 1)
/* add variables (if none are declared) */
(void) stata("restore", /* pre-save result */ 1)
(void) stata("append using 'tmpf'", 1)
(void) stata("save AP____DATA_pre-save.dta, replace", 1)
return(1)
}

'RM' appinv_runiform('RS' r, 'RS' c) return(runiform(r,c)) /* dummy */
'RM' appinv_rnormal('RS' r, 'RS' c) return(rnormal(r,c,0,1))

for(i=1;i<=rows((tmp=d#J(rows(d),1,1),J(rows(d),1,1)#d));i++) { /* start */
    (void) appinv_dist(tmp[i,1], tmp[i,2], n, /* zerod */ 0, /* normal */ 1)
    (void) appinv_dist(tmp[i,1], tmp[i,2], n, /* zerod */ 1, /* normal */ 1)
}

end
rm "AP____DATA_pre-save.dta" // delete pre-save
sca 't' = clock(c(current_date) + " " + c(current_time), "DMYhms") - 't'
di as res "Completed in '=round('t'/1000, .001)'s" // print timestamp

*** Bookmark #7
* ===== FINALIZE DATA ===== *
// drop obsolete
drop if r == 1 & c == 1

// order and sort
order variable N r c zerod nvar
sort r c zerod

// add data and variable labels
la data "Monte Carlo: mean NRMSE in general/zero-diagonal TM models, N = 50,000"
****
la var variable "variable name"
la var N "observations"
la var r "rows"
la var c "columns"
la var zerod "0 general TM, 1 zero-diagonal TM"
la var nvar "variables in 'a'"
la var rank "rank of 'a'"

// add data and variable notes
notes: Generated with random-number seed 123456789, descriptive /*
    */statistics collected with the help of summarizeby (SSC)

```

```
notes: Compiled by $compiled_by on TS

// save
compress                                // compress data
datasignature set, reset
****
save 'DTA_FILE', replace                // cloud-proof saving
copy 'DTA_FILE' "AP____DATA_$date.dta", replace
****
preserve
keep r c zerod mean sd p*
mata: fputmatrix((tmp=fopen("AP____DATA_$date.mmat", "w")), st_data(.))
mata: fclose(tmp)

***# Bookmark #8
* ===== CLEAN UP ===== *
cap log close data                      // close log
clear all
```

References

1. Nocedal, J.; Wright, S.J. *Numerical Optimization*, 2 ed.; Springer Series in Operations Research, Springer: New York, NY, 2006; p. 664. <https://doi.org/10.1007/978-0-387-40065-5>.

2. Boyd, S.; Vandenberghe, L. *Convex Optimization*, 1 ed.; Cambridge University Press: Cambridge, 2004; p. 727. <https://doi.org/10.1017/CBO9780511804441>.

3. Sydsæter, K.; Hammond, P.; Seierstad, A.; Strøm, A. *Further Mathematics for Economic Analysis*, 2 ed.; FT Prentice Hall: Harlow; München, 2011; p. 616.

4. Gentle, J.E. *Matrix Algebra: Theory, Computations and Applications in Statistics*, 3 ed.; Springer Texts in Statistics, Springer: Cham, 2024; p. 725. <https://doi.org/10.1007/978-3-031-42144-0>.

5. Dantzig, G.B. Reminiscences about the Origins of Linear Programming. *Memoirs of the American Mathematical Society* **1984**, *48*, 1–11. <https://doi.org/10.1090/memo/0298>.

6. Dantzig, G.B. *Linear Programming and Extensions*, 1 (reprint of 1963) ed.; Princeton Landmarks in Mathematics and Physics, Princeton University Press: Princeton, NJ, 1998; p. 656.

7. Koopmans, T.C. *Activity Analysis of Production and Allocation: Proceedings of a Conference*, 1 ed.; Wiley: New York, NY, 1951; p. 404.

8. Allen, R.G.D. *Mathematical Economics*, 2 (reprint of 1959) ed.; Palgrave Macmillan: London, 1976; p. 812. <https://doi.org/10.1007/978-1-349-81547-0>.

9. Lancaster, K. *Mathematical Economics*, 1 ed.; Dover Publications: New York, NY, 1987; p. 411.

10. Intriligator, M.D. *Mathematical Optimization and Economic Theory*, 1 (reprint of 1971) ed.; Classics in Applied Mathematics, SIAM: Philadelphia, PA, 2002; p. 508. <https://doi.org/10.1137/1.9780898719215>.

11. Intriligator, M.D.; Arrow, K.J. *Handbook of Mathematical Economics*, 1 ed.; Handbooks in Economics, North-Holland: Amsterdam; New York, NY, 1981; p. 378.

12. Dorfman, R.; Samuelson, P.A.; Solow, R.M. *Linear Programming and Economic Analysis*, 1 (reprint of 1958) ed.; Dover Books on Advanced Mathematics, Dover Publications: New York, NY, 1987; p. 525.

13. Frühwirth, T.; Abdennadher, S. *Essentials of Constraint Programming*, 1 ed.; Cognitive Technologies, Springer: Berlin; Heidelberg, 2003; p. 144. <https://doi.org/10.1007/978-3-662-05138-2>.

14. Rossi, F.; van Beek, P.; Walsh, T., Eds. *Handbook of Constraint Programming*, 1 ed.; Vol. 2, *Foundations of Artificial Intelligence*, Elsevier: Amsterdam; Boston, MA, 2006; p. 955.

15. Bolotov, I. Modeling of Time Series Cyclical Component on a Defined Set of Stationary Points and Its Application on the U.S. Business Cycle. In Proceedings of the The 8th International Days of Statistics and Economics. Melandrium, Sep 11–13 2014, pp. 151–160.

16. Bolotov, I. Modeling Bilateral Flows in Economics by Means of Exact Mathematical Methods. In Proceedings of the The 9th International Days of Statistics and Economics. Melandrium, Sep 10–12 2015, pp. 199–208.

17. Rosenthal, A. The History of Calculus. *The American Mathematical Monthly* **1951**, *58*, 75–86. <https://doi.org/10.2307/2308368>.

18. Stigler, S.M. *The History of Statistics: The Measurement of Uncertainty Before 1900*, 1 (reprint of 1986) ed.; Harvard University Press: Cambridge, MA, 2003; p. 410.

19. Ok, E.A. *Real Analysis with Economic Applications*, 1 ed.; Princeton University Press: Princeton, NJ, 2007; p. 802.
20. Creedy, J. The Early Use of Lagrange Multipliers in Economics. *The Economic Journal* **1980**, *90*, 371–376. <https://doi.org/10.2307/2231795>.
21. Plackett, R.L. Studies in the History of Probability and Statistics. XXIX: The Discovery of the Method of Least Squares. *Biometrika* **1972**, *59*, 239. <https://doi.org/10.2307/2334569>.
22. Stigler, S.M. Gauss and the Invention of Least Squares. *Annals of Statistics* **1981**, *9*, 465–474. <https://doi.org/10.1214/aos/1176345451>.
23. Stigler, S.M. *Statistics on the Table: The History of Statistical Concepts and Methods*, 1 (reprint of 1999) ed.; Harvard University Press: Cambridge, MA, 2002; p. 488.
24. Aitken, A.C. IV.—On Least Squares and Linear Combination of Observations. *Proceedings of the Royal Society of Edinburgh* **1936**, *55*, 42–48. <https://doi.org/10.1017/S0370164600014346>.
25. Aldrich, J. Doing Least Squares: Perspectives from Gauss and Yule. *International Statistical Review* **1998**, *66*, 61–81. <https://doi.org/10.2307/1403657>.
26. Mills, T.C.; Patterson, K., Eds. *Palgrave Handbook of Econometrics: Volume 1: Econometric Theory*, 1 ed.; Palgrave Macmillan: Basingstoke; New York, NY, 2006; p. 1097.
27. Mills, T.C.; Patterson, K., Eds. *Palgrave Handbook of Econometrics: Volume 2: Applied Econometrics*, 1 ed.; Palgrave Macmillan: Basingstoke; New York, NY, 2009; p. 1385.
28. Golub, G.H.; Kahan, W. Calculating the Singular Values and Pseudo-Inverse of a Matrix. *SIAM Journal on Numerical Analysis* **1965**, *2*, 205–224. <https://doi.org/10.1137/0702016>.
29. Eldén, L. A Weighted Pseudoinverse, Generalized Singular-Values, and Constrained Least-Squares Problems. *BIT Numerical Mathematics* **1982**, *22*, 487–502. <https://doi.org/10.1007/BF01934412>.
30. Ward, J.F. Limit Formula for Weighted Pseudoinverses. *SIAM Journal on Applied Mathematics* **1977**, *33*, 34–38. <https://doi.org/10.1137/0133003>.
31. Barata, J.C.A.; Hussein, M.S. The Moore-Penrose Pseudoinverse: A Tutorial Review of the Theory. *Brazilian Journal of Physics* **2012**, *42*, 146–165. <https://doi.org/10.1007/s13538-011-0052-z>.
32. Tikhonov, A.N.; Goncharskiy, A.V.; Stepanov, V.V.; Yagola, A.G. *Chislennyye metody resheniya nekorrektnykh zadach*, 2 ed.; Nauka: Moscow, 1990; p. 232.
33. Rao, C.R.; Mitra, S.K. *Generalized Inverse of Matrices and Its Applications*, 1 ed.; Wiley Series in Probability and Mathematical Statistics, Wiley: New York, NY, 1971; p. 240. <https://doi.org/10.1002/SERIES1345>.
34. Ben-Israel, A.; Greville, T.N.E. *Generalized Inverses: Theory and Applications*, 2 (reprint of 2003) ed.; CMS Books in Mathematics, Springer: New York, NY, 2006; p. 420. <https://doi.org/10.1007/b97366>.
35. Dresden, A. The Fourteenth Western Meeting of the American Mathematical Society. *Bulletin of the American Mathematical Society* **1920**, *26*, 385–397. <https://doi.org/10.1090/S0002-9904-1920-03322-7>.
36. Bjerhammar, A. Rectangular Reciprocal Matrices, With Special Reference to Geodetic Calculations. *Bulletin géodésique* **1951**, *20*, 188–220. <https://doi.org/10.1007/BF02526278>.
37. Penrose, R. A Generalized Inverse for Matrices. *Mathematical Proceedings of the Cambridge Philosophical Society* **1955**, *51*, 406–413. <https://doi.org/10.1017/S0305004100030401>.
38. Penrose, R. On Best Approximate Solutions of Linear Matrix Equations. *Mathematical Proceedings of the Cambridge Philosophical Society* **1956**, *52*, 17–19. <https://doi.org/10.1017/S0305004100030929>.
39. Chipman, J.S. On Least Squares with Insufficient Observations. *Journal of the American Statistical Association* **1964**, *59*, 1078–1111. <https://doi.org/10.2307/2282625>.
40. Price, C.M. The Matrix Pseudoinverse and Minimal Variance Estimates. *SIAM Review* **1964**, *6*, 115–120. <https://doi.org/10.1137/1006029>.
41. Plackett, R.L. Some Theorems in Least Squares. *Biometrika* **1950**, *37*, 149–157. <https://doi.org/10.2307/2332158>.
42. Rao, C.R.; Mitra, S.K., Generalized Inverse of a Matrix and Its Applications. In *Theory of Statistics: Proceedings of the Sixth Berkeley Symposium on Mathematical Statistics and Probability, Volume I*, 1 ed.; Le Cam, L.M.; Neyman, J.; Scott, E.L., Eds.; University of California Press: Berkeley, CA, 1972; pp. 601–620.
43. Greville, T.N.E. The Pseudoinverse of a Rectangular Matrix and Its Statistical Applications. In Proceedings of the Annual Meeting of the American Statistical Association. American Statistical Association, Dec 27–30 1958, pp. 116–121.
44. Greville, T.N.E. The Pseudoinverse of a Rectangular or Singular Matrix and Its Application to the Solution of Systems of Linear Equations. *SIAM Review* **1959**, *1*, 38–43. <https://doi.org/10.1137/1001003>.

45. Greville, T.N.E. Some Applications of the Pseudoinverse of a Matrix. *SIAM Review* **1960**, *2*, 15–22. <https://doi.org/10.1137/1002004>.
46. Greville, T.N.E. Note on Fitting of Functions of Several Independent Variables. *Journal of the Society for Industrial and Applied Mathematics* **1961**, *9*, 109–115. <https://doi.org/10.1137/0109011>.
47. Cline, R.E. Representations for the Generalized Inverse of a Partitioned Matrix. *Journal of the Society for Industrial and Applied Mathematics* **1964**, *12*, 588–600. <https://doi.org/10.1137/0112050>.
48. Cline, R.E.; Greville, T.N.E. An Extension of the Generalized Inverse of a Matrix. *SIAM Journal on Applied Mathematics* **1970**, *19*, 682–688. <https://doi.org/10.1137/0119069>.
49. Ben-Israel, A.; Cohen, D. On Iterative Computation of Generalized Inverses and Associated Projections. *SIAM Review* **1966**, *8*, 410–419. <https://doi.org/10.1137/0703035>.
50. Lewis, T.O.; Newman, T.G. Pseudoinverses of Positive Semidefinite Matrices. *SIAM Journal on Applied Mathematics* **1968**, *16*, 701–703. <https://doi.org/10.1137/0116057>.
51. Bott, R.; Duffin, R.J. On the Algebra of Networks. *Transactions of the American Mathematical Society* **1953**, *74*, 99–109. <https://doi.org/10.2307/1990850>.
52. Campbell, S.L.; Meyer, C.D. *Generalized Inverses of Linear Transformations*, 1 (reprint of 1979) ed.; Classics in Applied Mathematics, SIAM: Philadelphia, PA, 2009; p. 272. <https://doi.org/10.1137/1.9780898719048>.
53. Wang, G.; Wei, Y.; Qiao, S. *Generalized Inverses: Theory and Computations*, 1 ed.; Vol. 53, *Developments in Mathematics*, Springer: Singapore, 2018; p. 397. <https://doi.org/10.1007/978-981-13-0146-9>.
54. Meyer, Carl D., J. Generalized Inverses and Ranks of Block Matrices. *SIAM Journal on Applied Mathematics* **1973**, *25*, 597–602. <https://doi.org/10.1137/0125057>.
55. Hartwig, R.E. Block Generalized Inverses. *Archive for Rational Mechanics and Analysis* **1976**, *61*, 197–251. <https://doi.org/10.1007/bf00281485>.
56. Rao, C.R.; Yanai, H. Generalized Inverses of Partitioned Matrices Useful in Statistical Applications. *Linear Algebra and Its Applications* **1985**, *70*, 105–113. [https://doi.org/10.1016/0024-3795\(85\)90046-1](https://doi.org/10.1016/0024-3795(85)90046-1).
57. Tian, Y. The Moore-Penrose Inverses of $m \times n$ Block Matrices and Their Applications. *Linear Algebra and Its Applications* **1998**, *283*, 35–60. [https://doi.org/10.1016/s0024-3795\(98\)10049-6](https://doi.org/10.1016/s0024-3795(98)10049-6).
58. Rakha, M.A. On the Moore-Penrose Generalized Inverse Matrix. *Applied Mathematics and Computation* **2004**, *158*, 185–200. <https://doi.org/10.1016/j.amc.2003.09.004>.
59. Baksalary, O.M.; Trenkler, G. On Formulae for the Moore-Penrose Inverse of a Columnwise Partitioned Matrix. *Applied Mathematics and Computation* **2021**, *403*, 1–10. <https://doi.org/10.1016/j.amc.2020.125913>.
60. Albert, A.E. *Regression and the Moore-Penrose Pseudoinverse*, 1 ed.; Mathematics in Science and Engineering, Academic Press: New York, NY, 1972; p. 180.
61. Dokmanić, I.; Kolundžija, M.; Vetterli, M. Beyond Moore-Penrose: Sparse Pseudoinverse. In Proceedings of the IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP). IEEE, May 26–31 2013, pp. 6526–6530. <https://doi.org/10.1109/ICASSP.2013.6638923>.
62. Baksalary, O.M.; Trenkler, G. The Moore-Penrose Inverse: A Hundred Years on a Frontline of Physics Research. *European Physical Journal H* **2021**, *46*, 1–10. <https://doi.org/10.1140/epjh/s13129-021-00011-y>.
63. Mortari, D. Least-Squares Solution of Linear Differential Equations. *Mathematics* **2017**, *5*, 48. <https://doi.org/10.3390/math5040048>.
64. Lawson, C.L.; Hanson, R.J. *Solving Least Squares Problems*, 1 (reprint of 1974) ed.; Classics in Applied Mathematics, SIAM: Philadelphia, PA, 1995; p. 337. <https://doi.org/10.1137/1.9781611971217>.
65. Getson, A.J.; Hsuan, F.C. *{2}-Inverses and Their Statistical Application*, 1 (reprint of 1988) ed.; Lecture Notes in Statistics, Springer: New York, NY, 2012; p. 110. <https://doi.org/10.1007/978-1-4612-3930-7>.
66. Björck, A. *Numerical Methods for Least Squares Problems*, 1 ed.; SIAM: Philadelphia, PA, 1996; p. 408.
67. Kantorovich, L.V. *Matematicheskiye metody organizatsii i planirovaniia proizvodstva*, reprint ed.; Izdatel'skiy dom S.-Peterb. gos. un-ta: Saint Petersburg, 2012; p. 96.
68. Shamir, R. The Efficiency of the Simplex Method: A Survey. *Management Science* **1987**, *33*, 301–334. <https://doi.org/10.1287/mnsc.33.3.301>.
69. Stone, R.E.; Tovey, C.A. The Simplex and Projective Scaling Algorithms as Iteratively Reweighted Least Squares Methods. *SIAM Review* **1991**, *33*, 220–237. <https://doi.org/10.1137/1033049>.
70. Wagner, H.M. Linear Programming Techniques for Regression Analysis. *Journal of the American Statistical Association* **1959**, *54*, 206–212. <https://doi.org/10.1080/01621459.1959.10501506>.
71. Sielken, R.L.; Hartley, H.O. Two Linear Programming Algorithms for Unbiased Estimation of Linear Models. *Journal of the American Statistical Association* **1973**, *68*, 639–641. <https://doi.org/10.1080/01621459.1973.10481398>.

72. Kiountouzis, E.A. Linear Programming Techniques in Regression Analysis. *Applied Statistics* **1973**, *22*, 69. <https://doi.org/10.2307/2346304>.
73. Spósito, V.A. On Unbiased Lp Regression Estimators. *Journal of the American Statistical Association* **1982**, *77*, 652–653. <https://doi.org/10.1080/01621459.1982.10477865>.
74. Judge, G.G.; Takayama, T. Inequality Restrictions in Regression Analysis. *Journal of the American Statistical Association* **1966**, *61*, 166–181. <https://doi.org/10.2307/2283052>.
75. Mantel, N. Restricted Least Squares Regression and Convex Quadratic Programming. *Technometrics* **1969**, *11*, 763–773. <https://doi.org/10.2307/1266897>.
76. Blair, C. Random Linear Programs with Many Variables and Few Constraints. *Mathematical Programming* **1986**, *34*, 62–71. <https://doi.org/10.1007/bf01582163>.
77. Donoho, D.L.; Tanner, J. Sparse Nonnegative Solution of Underdetermined Linear Equations by Linear Programming. *Proceedings of the National Academy of Sciences of the United States of America* **2005**, *102*, 9446–9451. <https://doi.org/10.1073/pnas.0502269102>.
78. Osborne, M.R. Degeneracy: Resolve or Avoid? *Journal of the Operational Research Society* **1992**, *43*, 829–835. <https://doi.org/10.2307/2583102>.
79. George, K.; Osborne, M.R. On Degeneracy in Linear Programming and Related Problems. *Annals of Operations Research* **1993**, *46–47*, 343–359. <https://doi.org/10.1007/BF02023104>.
80. Stoer, J. On the Numerical Solution of Constrained Least-Squares Problems. *SIAM Journal on Numerical Analysis* **1971**, *8*, 382–411. <https://doi.org/10.1137/0708038>.
81. Waterman, M.S. A Restricted Least Squares Problem. *Technometrics* **1974**, *16*, 135–136. <https://doi.org/10.1080/00401706.1974.10489160>.
82. Grafarend, E.W.; Awange, J.L., Algebraic Solutions of Systems of Equations. In *Linear and Nonlinear Models*, 1 ed.; Springer: Berlin; Heidelberg, 2012; pp. 527–569. https://doi.org/10.1007/978-3-642-22241-2_15.
83. Qian, J.; Andrew, A.L.; Chu, D.; Tan, R.C.E. Methods for Solving Underdetermined Systems. *Numerical Linear Algebra with Applications* **2018**, *25*, 17. <https://doi.org/10.1002/nla.2127>.
84. Wolfe, P. A Technique for Resolving Degeneracy in Linear Programming. *Journal of the Society for Industrial and Applied Mathematics* **1963**, *11*, 205–211. <https://doi.org/10.1137/0111016>.
85. Dax, A. Linear Programming via Least Squares. *Linear Algebra and Its Applications* **1988**, *111*, 313–324. [https://doi.org/10.1016/0024-3795\(88\)90067-5](https://doi.org/10.1016/0024-3795(88)90067-5).
86. Barnes, E.; Chen, V.; Gopalakrishnan, B.; Johnson, E.L. A Least-Squares Primal-Dual Algorithm for Solving Linear Programming Problems. *Operations Research Letters* **2002**, *30*, 289–294. [https://doi.org/10.1016/s0167-6377\(02\)00163-3](https://doi.org/10.1016/s0167-6377(02)00163-3).
87. Bei, X.; Chen, N.; Zhang, S. Solving Linear Programming with Constraints Unknown. In *Proceedings of the 42nd International Colloquium, ICALP 2015*. Springer, Jul 6–10 2015, Vol. 9134, pp. 129–142. https://doi.org/10.1007/978-3-662-47672-7_16.
88. Übi, E. Exact and Stable Least Squares Solution to the Linear Programming Problem. *Central European Journal of Mathematics* **2005**, *3*, 228–241. <https://doi.org/10.2478/BF02479198>.
89. Übi, E. On Stable Least Squares Solution to the System of Linear Inequalities. *Open Mathematics* **2007**, *5*, 373–385. <https://doi.org/10.2478/s11533-007-0003-7>.
90. Übi, E. A Numerically Stable Least Squares Solution to the Quadratic Programming Problem. *Open Mathematics* **2008**, *6*, 171–178. <https://doi.org/10.2478/s11533-008-0012-1>.
91. Übi, E. Mathematical Programming via the Least-Squares Method. *Open Mathematics* **2010**, *8*. <https://doi.org/10.2478/s11533-010-0049-9>.
92. Übi, E. Linear Inequalities via Least Squares. *Proceedings of the Estonian Academy of Sciences* **2013**, *62*, 238–248. <https://doi.org/10.3176/proc.2013.4.04>.
93. Whiteside, M.M.; Choi, B.; Eakin, M.; Crockett, H. Stability of Linear Programming Solutions Using Regression Coefficients. *Journal of Statistical Computation and Simulation* **1994**, *50*, 131–146. <https://doi.org/10.1080/00949659408811605>.
94. Blair, C. Random Inequality Constraint Systems with Few Variables. *Mathematical Programming* **1986**, *35*, 135–139. <https://doi.org/10.1007/bf01580644>.
95. Pereira-López, X.; Fernández-Fernández, M.; Carrascal-Incera, A. Rectangular Input-output Models by Moore-Penrose Inverse. *Revista Electrónica De Comunicaciones Y Trabajos De ASEPUMA* **2014**, *15*, 13–24.
96. Bolotov, I. The Problem of Relationships Between Conditional Statements and Arithmetic Functions. *Mundus Symbolicus* **2012**, *20*, 5–12.

97. Bolotov, I. SUMMARIZEBY: Stata Module to Use Statsby Functionality With Summarize (version 1.1.2), 2022.
98. Bolotov, I. PYCLSP: Modular Two-Step Convex Optimization Estimator for Ill-Posed Problems (version 1.3.0), 2025.
99. Bolotov, I. PYTMPINV: Tabular Matrix Problems via Pseudoinverse Estimation (version 1.2.0), 2025.
100. Bolotov, I. PYLPPINV: Linear Programming via Pseudoinverse Estimation (version 1.3.0), 2025.

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.