

Aşırı Derecede Küçük Örneklem Problemleri için Hibrit Regresyon Modellerde Maksimum Entropi Kovaryans Tahmin Edicisi

Maximum Entropy Covariance Estimator in Hybrid Regression Models for Extremely Small Sample Problems

Mamadou Diallo^{1*}

 ÖZET

Aşır Genç²



¹ Necmettin Erbakan Üniversitesi,
Fen Fakültesi, İstatistik Bölümü,
Konya, Türkiye

² Necmettin Erbakan Üniversitesi,
Fen Fakültesi, İstatistik Bölümü,
Konya, Türkiye

*Sorumlu Yazar/Corresponding
Author

Bu çalışma, gözlem sayısının değişken sayısına göre son derece küçük olduğu ($n \ll p$) durumlarda kovaryans matrisinin kararsızlığı nedeniyle ortaya çıkan tahmin sorunlarını ele alarak Hibrit Regresyon Modelinin (HRM) performansını incelemektedir. Çalışmanın yöntemi iki bileşen üzerine kuruludur: (i) Hibrit Kovaryans Tahmin Edicisi (HCE) ile spektral dengeleme yoluyla kovaryans matrisinin stabilize edilmesi ve (ii) farklı regresyon yapılarının (HRM1–HRM5) optimal ağırlıklarla birleştirildiği hibritleştirilmiş yaklaşımı. HCE, maksimum entropi tabanlı kovaryans tahminini çeşitli düzenlenleştirilmiş yapılarla bütünleştirerek kötü koşullanmış matrislerin sayısal güvenilirliğini artırmaktadır. Farklı örneklem büyüklükleri ve boyutsal senaryolar altında yürütülen simülasyonlar, stabilizasyon ve hibritleştirmeyi birlikte kullanan HRM4 modelinin AIC, CAIC ve ICOMP kriterlerine göre tutarlı biçimde en düşük değerlere ulaşarak en iyi performansı gösterdiğini ortaya koymaktadır. Sonuçlar, kovaryansın bilgi karmaşıklığının model seçimi sürecinde belirleyici olduğunu ve HCE tabanlı hibrit yaklaşımın küçük örneklem probleminde yalnlık–varyans dengesini etkili biçimde optimize ettiğini göstermektedir. Çalışma, yöntemin yüksek boyutlu veri ortamlarına genellenebilir olduğunu ve gelecekte diğer düzenlenleştirilmiş kovaryans tahmin edicileriyle karşılaştırmalı analizlerin yapılabileceğini vurgulamaktadır.

Anahtar Kelimeler: Maksimum Entropi, Hibrit kovaryans tahmincisi (HCE), Hibrit regresyon modeli, Küçük örneklem problemi

Cite this Article: “Mamadou, D.; Aşır, G. (2025). “Aşırı Derecede Küçük Örneklem Problemleri için Hibrit Regresyon Modellerde Maksimum Entropi Kovaryans Tahmin Edicisi”, *Journal of Economic and Social Analysis (JESA)*. (1), 20-38.

ABSTRACT

This study investigates the performance of the Hybrid Regression Model (HRM) in settings where the number of observations is extremely small compared to the number of variables ($n \ll p$), causing severe instability in covariance estimation. The methodological framework combines two complementary components: (i) spectral stabilization of the covariance matrix through the Hybrid Covariance Estimator (HCE), which integrates maximum-entropy-based estimation with several regularized covariance structures, and (ii) model hybridization, where multiple regression formulations (HRM1–HRM5) are combined using optimal weights to achieve an improved bias–variance trade-off. Simulation experiments conducted across a wide range of sample sizes and dimensional settings demonstrate that the HRM4 model, which jointly employs covariance stabilization and hybridization, consistently achieves the lowest AIC, CAIC, and ICOMP scores. This confirms its superior numerical stability, robustness, and predictive accuracy in undersized sample scenarios. The findings highlight the crucial role of covariance information complexity in model selection and show that HCE-based hybrid regression provides an effective strategy for controlling instability in high-dimensional regression. The study concludes by suggesting extensions of the approach to other regularized covariance estimators and applications to real high-dimensional datasets.

Keywords: Maximum Entropy, Hybrid Covariance Estimator (HCE), Hybrid Regression Model, Small Sample Problem

1 | GİRİŞ

Geleneksel istatistiksel analiz yöntemleri, genellikle gözlem sayısının değişken sayısından fazla olduğu varsayımına dayanmaktadır. Ancak modern veri analizinde, özellikle genetik, metin madenciliği ve finans gibi alanlarda, örneklem büyüklüğünün (n) değişken sayısına (p) göre çok küçük olduğu durumlar ($n \ll p$) sıkça karşılaşılmaktadır. Bu durum, örnek kovaryans matrisinin tekil veya kötü şartlandırılmış olmasına neden olarak klasik çok değişkenli analizlerin doğruluğunu ve geçerliliğini tehlikeye atmaktadır (Donoho, 2000; Stein, 1975). Bu tür durumlarda, incelenebilecek örnekler genellikle birkaç düzine veya yüzlerce kişiyle sınırlı olsa da tek bir gözlem, özellikle daha önce bahsettiğimiz alanlarda binlerce hatta milyonlarca değişkeni içerebilir. Ancak, klasik yöntemler bu tür verileri verimli bir şekilde işlemek üzere tasarlanmamıştır (Donoho, 2000). Ayrıca Stein'in (1956, 1975) uzun zaman önce ortaya koyduğu gibi, kovaryans yapısı Σ olan ve ortalaması sıfır kabul edilen normal dağılmış bir popülasyondan alınan n -boyutlu bir örneklem varyans-kovaryans matrisinin en çok olabilirlik tahmini, p/n oranı yüksek olduğunda güvenilir bir tahmin edici değildir. Bununla birlikte, söz konusu tahmin edici yanlış olmamakta ve pozitif tanımlı kalmaktadır (Stein, 1975). Bu durumda, kovaryans matrisinin yapısı bozulmakta ve büyük özdeğerler yukarı, küçük özdeğerler ise aşağı yönlü önyargılı olmaktadır (Yilun Chen vd., 2011). Buna karşılık, $p < n$ olduğunda özdeğerler daha hızlı azalır ve genel olarak daha düşük kalır; bu ise kovaryans matrisinin daha istikrarlı bir tahminini yansıtır. Bu önyargı, p/n oranının artmasıyla birlikte artmakta; yani değişken sayısının gözlem sayısına oranı ne kadar yüksekse, tahmin edilen kovaryans matrisi o kadar kararsız olmaktadır. Bu durum, kovaryans matrisi tahmin kalitesinde bozulmaya işaret eder ve aynı zamanda temel bileşen analizi gibi çok değişkenli analizlerin doğruluğunu olumsuz etkileyebilir. Bu bağlamda, literatürde çeşitli düzgünleştirilmiş (smoothed) veya büzülmüş (shrinkage) kovaryans tahmin yöntemleri önerilmiştir. Bir dizi çalışmanın ardından, yüksek boyutlu veri kümelerinde kovaryans problemine çözüm getirmek amacıyla çeşitli yaklaşımlar önerilmiştir. İstatistikçiler bu durumu sıklıkla “büyük p , küçük n ” ifadesiyle tanımlar; bu, sınırlı sayıda gözleme kıyasla çok fazla değişken bulunduğunu ifade eder. Bu bağlamda, Hibrit Kovaryans Tahmin Edicisi (Hybrid Covariance Estimator, HCE) literatüre ilk kez Pamukcu ve ark. (2015) tarafından kazandırılmıştır (Pamukçu vd., 2015). Bu tahmin edici, maksimum entropi kovaryans tahmin edicisinin, belirli yumuşatılmış kovaryans yapıları ile birleştirilmesine dayanmaktadır (Fiebig, 1984). HCE sayesinde kovaryans yapısındaki bozulma azaltılabilmekte, bu da özellikle $n \ll p$ durumuyla karakterize edilen yüksek boyutlu veri kümelerinde çok değişkenli istatistiksel yöntemlerin daha etkin

uygulanmasına olanak sağlamaktadır. Bu çalışmanın amacı, hem istikrar kazandırma hem de düzenleştirme tekniklerini birleştirerek ve ayrıca bilgi kriterlerini kullanarak, $n \ll p$ koşulu ile karakterize edilen yüksek boyutlu veri kümelerinin analizinde Hibrit Kovaryans Tahmin Edicisi'ne (HCE) dayalı regresyon analizini tanıtmaktır.

Daha spesifik olarak, bu çalışma şu şekilde yapılandırılmıştır: Bölüm 2.1'de Ridge Regresyonu ve düzleştirilmiş (smoothed) kovaryans tahmin edicileri sunulmaktadır; Bölüm 2.2'de Maksimum Entropi ve Hibrit Kovaryans Tahmini (HCE) tanımlanmaktadır ve Bölüm 2.3'te ICOMP ve HCE'ye Dayalı Model Seçim Kriterleri ve Hibrit Regresyon Modeli (HRM) ele alınmaktadır. Bölüm 3'te, Simülasyon protokolüne genel bir bakış açısıyla farklı p ve n değerleriyle tanımlanan çeşitli senaryolara ve γ simülasyon sayısı HRM modelinin uygulandığı bir simülasyon çalışmasını sunmaktadır. Son olarak, Bölüm 4, elde edilen sonuçların tartışılmasını, analizini ve Bölüm 5 ise gelecekteki perspektifleri sunmaktadır.

2 | Materyal Ve Metot

2.1 | Ridge Regresyonu ve Düzleştirilmiş (Smoothed) Kovaryans Tahmin Edicileri

Klasik doğrusal modellerde, En Küçük Kareler (Ordinary Least Squares – OLS) yöntemiyle yapılan tahmin, $X'X$ matrisinin tersinin alınabilir olması varsayımına dayanmaktadır. Ancak, açıklayıcı değişkenler arasında güçlü bir korelasyon bulunduğunda (çoklu doğrusal bağlantı -multicollinearity) ya da değişken sayısı p , gözlem sayısını n 'den fazla olduğunda (yani $n \ll p$ durumunda), veri matrisi kötü koşullu (ill-conditioned) hâle gelir veya tekil olmaktadır. Bu koşullar altında klasik yöntemler (OLS tahmin edicileri) çok yüksek varyansa, sayısal kararsızlığa ve zayıf tahmin performansına yol açmaktadır. Bu sınırlamaları gidermek için, kovaryans matrisini değiştirmeyi veya kararlı hâle getirmeyi amaçlayan çeşitli düzenleştirme teknikleri geliştirilmiştir; aşağıda bunlardan her birine kısaca değinilecektir.

Hoerl ve Kennard, ilk olarak ridge regresyon tahmin edicisini önermişlerdir (A. E. Hoerl & Kennard, 1970). Bu yöntem, doğrusal regresyon modellerinde çoklu doğrusal bağlantı (multicollinearity) sorununu çözmek için geliştirilen en etkili yaklaşımlardan biridir. Ridge yöntemi, OLS (En Küçük Kareler) maliyet fonksiyonuna eklenen bir ikinci dereceden (kuadratik) cezalandırma terimi aracılığıyla katsayıların büyüklüğünü sınırlandırır ve tahminin kararlılığını artırmaktadır. İlgili optimizasyon problemi şu şekilde ifade edilmektedir:

$$\hat{\beta}_{ridge} = \arg \min \{\|y - X\beta\|^2 + \|\alpha\beta\|^2\} \quad (1)$$

Ridge regresyon tahmin edicisi, (1) ifadesinin β 'ya göre minimize ederek elde edilmektedir. Buna göre Ridge kovaryans matrisi,

$$\hat{\Sigma}_R = \hat{\Sigma}_{MLE} + \alpha I_p \quad (2)$$

şeklinde hesaplanmaktadır.

Burada αI_p terimi, XX' matrisinin özdeğer spektrumunu değiştirerek sayısal hataların büyümesini önler. Uygulamada, küçük bir α değeri çözümü klasik yöntemlere (OLS) yaklaştırırken, büyük bir α değeri katsayıları sıfıra doğru sınırlar. Sapma ile varyans arasındaki bu denge, modelin kararlılığını ve genellenebilirlik kapasitesini artırmaktadır (R. W. Hoerl, 2020). Bu fikirden yola çıkarak farklı düzenleme biçimleri geliştirilmiştir. Tibshirani (1996), “en küçük mutlak değer daraltma ve seçim işleci” anlamına gelen yeni bir teknik, yani lasso yöntemini önermiştir (Tibshirani, 1996). Bu yaklaşım, L_2 normu yerine L_1 cezası kullanır; böylece bazı katsayıların tam olarak sıfır olmasına yol açar ve değişkenlerin otomatik olarak seçilmesini sağlar. Sürekli bir daraltma yöntemi olarak ridge regresyon, ön yargı–varyans dengesi sayesinde en iyi tahmin performansına ulaşmaktadır (Tibshirani, 1996). Ancak ridge regresyon, tüm yordayıcıları (değişkenleri) modelde tuttuğundan dolayı sade (parsimonious) bir model üretememektedir. Buna karşılık, $p > n$ durumunda lasso yöntemi, dışbükey optimizasyon probleminin yapısı gereği en fazla n değişken seçebilmekte; bu sınır aşıldığında ise değişken seçiminde doyuma ulaşmaktadır. Zou ve Hastie, 2005’te, iki cezalandırmayı (Lasso ve Ridge) birleştirerek Ridge’in sayısal kararlılığını Lasso’nun sadeliğiyle uzlaştırır, özellikle yüksek korelasyonlu değişkenler durumunda etkili bir çözüm sunmaktadır (Zou & Hastie, 2005).

Buna paralel olarak, kovaryans matrisinin düzenlenilmesi meselesi, yüksek boyutlu çok değişkenli analizlerde merkezî bir araştırma konusu hâline gelmiştir. Klasik ampirik tahmin edici:

$$\hat{\Sigma}_{MLE} = \frac{1}{n-1} (X - \bar{X})(X - \bar{X})' \quad (3)$$

n, p 'ye kıyasla küçük olduğunda tekil (kararsız) hâle gelmektedir. Bu sorunu aşmak için Ledoit ve Wolf (2004), deneysel matris $\hat{\Sigma}_{MLE}$ ile iyi koşullandırılmış hedef matris T arasında ağırlıklı bir kombinasyon olarak tanımlanan bir büzülme tahmincisi önermiştir (Ledoit & Wolf, 2004).

$$\hat{\Sigma}_{LW} = (1 - \delta)\hat{\Sigma}_{MLE} + \delta T \quad (4.a)$$

Burada δ , düzenleme parametresidir (büzülme katsayısı) ve 0 ile 1 arasında bir değer alır. Bu değer, gözlemlerin bir fonksiyonu da olabilir. T matrisi, büzülme hedefi olarak adlandırılmaktadır. T 'nin basit hali şöyledir:

$$\hat{T} = \frac{tr(\hat{\Sigma}_{MLE})}{p} I_p = \left(\frac{1}{p} \sum_{j=1}^p \lambda_j \right) I_p = \bar{\lambda} I_p \quad (4.b)$$

Burada tr , $\hat{\Sigma}_{MLE}$ 'nin köşegeninin değerlerinin toplamıdır, λ_j $j=1, \dots, p$ tahmini örnek kovaryans matrisinin özdeğerleridir, ve $\bar{\lambda}$ özdeğerlerin aritmetik ortalamasıdır. Bu tür tahmin (4.a), kontrollü bir ön yargı pahasına varyansı azaltmaya olanak tanır ve böylece sonlu örneklemede daha iyi bir performans sağlamaktadır.

Hibrit Kovaryans Tahmincisi (HCE) gibi daha yeni yaklaşımlar, $n \ll p$ 'nin olduğu aşırı durumlarda bile kovaryans matrisinin sağlam bir tahminini elde etmek için stabilizasyon ve düzenleme prensiplerini birleştirmektedir. Bu çalışmada, bu sorunu çözmek için çeşitli düzenlenmiş veya düzeltilmiş kovaryans tahmincileri kullanılmıştır. Aşağıda listelenen bu tahminciler, analizimizin metodolojik temelini oluşturmaktadır.

Bozdoğan'ın Convex Sum tahmin edicisi (Bozdoğan, 2009).

$$\hat{\Sigma}_{BCSE} = \hat{p} \hat{\Sigma} + (1 - \hat{p}) \hat{D} \quad (5.a)$$

$$\hat{p} = \frac{1}{\alpha} \quad (5.b)$$

$$\alpha = \left(\frac{1}{n-1} \sum_{j=1}^p Var(x_j) \right) \quad (5.c)$$

Empirical Bayes tahmin edicisi (Haff, 1980).

$$\hat{\Sigma}_{EB} = \hat{\Sigma} + \frac{(p-1)}{n \cdot tr(\hat{\Sigma})} I_p \quad (6)$$

Convex Sum tahmin edicisi (Press, 1975).

$$\hat{\Sigma}_{CSE} = + \frac{n}{m+n} \hat{\Sigma} + \left(1 - \frac{n}{m+n} \right) \left[\frac{tr(\hat{\Sigma})}{p} \right] I_p \quad (7.a)$$

$$0 < m < \frac{2[p(1+\beta)-2]}{p-\beta} \quad (7.b)$$

$$\beta = \frac{tr(\hat{\Sigma})^2}{tr(\hat{\Sigma}^2)} \quad (7.c)$$

Stipulated Ridge tahmin edicisi (Shurygin, 1983).

$$\hat{\Sigma}_{SRE} = \hat{\Sigma} + p(p-1) [2n \cdot tr(\hat{\Sigma})]^{-1} I_p \quad (8)$$

Stipulated Diagonal tahmin edicisi (Shurygin, 1983).

$$\hat{\Sigma}_{SDE} = (1 - \pi) \hat{\Sigma} + \pi \text{diag}(\hat{\Sigma}) \quad (9.a)$$

$$\pi = p(p-1) [2n \cdot tr(\hat{\Sigma}^{-1}) - p]^{-1} \quad (9.b)$$

2.2 | Maksimum Entropi ve Hibrit Kovaryans Tahmini (HCE)

Birçok istatistiksel uygulamada, özellikle çok değişkenli analiz bağlamında, kovaryans matrislerinin doğru şekilde tahmin edilmesi verinin temel yapısını anlamada kritik bir öneme sahiptir. Ancak klasik maksimum olabilirlik (MLE) temelli kovaryans tahmin edicileri,

değişken sayısının örneklem büyüklüğünü aştığı durumlarda kararsız hatta tekil hâle gelmekte ve bu durum söz konusu yöntemlerin kullanılabilirliğini ciddi biçimde kısıtlamaktadır. “Boyutsallık laneti” olarak adlandırılan bu problem, hem sağlam hem de sayısal açıdan kararlı kovaryans tahmin yöntemlerinin geliştirilmesini zorunlu kılmaktadır (Pamukçu vd., 2015). Alternatif yaklaşımlar arasında, Maksimum Entropi’ye (Maximum Entropy Covariance Estimator - MCE) dayalı yöntemler, maksimum belirsizlik ilkesine dayanarak kovaryansın daha düzgün bir tahminini elde etmeye olanak tanır; bu, özdeğerler üzerine belirli kısıtlamalar getirilerek gerçekleştirilmiştir (Jaynes, 1957). Özellikle özdeğerlerin çoğunun negatif veya sıfır olduğu durumlar için uygun bir çözüm, verilerin temel yapısını korurken tahmini kararlı hâle getiren maksimum entropili kovaryans matrisinin kullanılmasıdır.

2.2.1 | Veri matrisi formülasyonu ve düzenleme yoluyla stabilizasyon

Basitleştirilmiş durumda, iki değişken x ve y ’nin olduğu çok değişkenli gözlemlerden oluşan bir $[Z]$ matrisi

$$Z = \begin{bmatrix} x_1 & y_1 \\ \vdots & \vdots \\ x_n & y_n \end{bmatrix} \quad (10)$$

dir. $[Z]$ ’den tahmin edilen kovaryans matrisi, örneklemdeki bitişik gözlemler arasındaki farklardan ceza terimlerini hesaplayarak kovaryansı sabitlemeye yarayan bir köşegen matrisi $[D]$ eklenerek düzenlenebilmektedir. Bu matris şu şekilde tanımlanmaktadır:

$$D = \begin{bmatrix} \frac{1}{12n} \sum (\xi_j - \xi_{j-1})^2 + \frac{1}{6n} \sum \{ (\xi_1 + \xi_0)^2 (\xi_n + \xi_{n-1})^2 \} & 0 \\ 0 & \frac{1}{12n} \sum (\eta_j - \eta_{j-1})^2 + \frac{1}{6n} \sum \{ (\eta_1 - \eta_0)^2 (\eta_n - \eta_{n-1})^2 \} \end{bmatrix} \quad (11)$$

burada ξ_j ve η_j birincil ortalama noktalarıdır:

$$\xi_j = \frac{1}{2} (x^j + x^{j+1}) \quad j=0,1,\dots,n \quad (12.a)$$

$$\eta_j = \frac{1}{2} (y^j + y^{j+1}) \quad j=0,1,\dots,n \quad (12.b)$$

$$\text{ve } x^0 \equiv x^1 < x^2 < \dots x^n \equiv x^{n+1} \quad (13.a)$$

$$y^0 \equiv y^1 < y^2 < \dots y^n \equiv y^{n+1} \quad (13.b)$$

Sıra istatistikleridir, düzenlenilmenin dayandığı karakteristik noktaları ifade etmektedir. Bu düzenleme, sayısal kararlılığı artırır ve yüksek boyutlu matrislerde aşırı uyum (overfitting) etkilerini sınırlamaya olanak tanımaktadır. (12.a) ve (12.b) numaralı denklemde kullanılan ikincil ortalama noktaları ise şu şekilde hesaplanmaktadır:

$$\bar{x}^j = \frac{1}{2} (\xi_{j-1} + \xi_j) \quad j=0,1,\dots,n \quad (14.a)$$

$$\bar{y}^j = \frac{1}{2}(\eta_{j-1} + \eta_j) \quad (14.b)$$

Daha sonra C matrisi, bu değiştirilmiş değerlerin ampirik kovaryans matrisi,

$$C = \begin{bmatrix} \frac{1}{n} \sum (\bar{x}_k - \bar{x})^2 & \frac{1}{n} \sum (\bar{y}_k - \bar{y})(\bar{x}_k - \bar{x}) \\ \frac{1}{n} \sum (\bar{y}_k - \bar{y})(\bar{x}_k - \bar{x}) & \frac{1}{n} \sum (\bar{y}_k - \bar{y})^2 \end{bmatrix} \quad (15)$$

olarak hesaplanmaktadır. Burada $\bar{x}$ ve $\bar{y}$, $\bar{x}_k$ ve $\bar{y}_k$ 'nin ilgili ortalamalarıdır. Bu yapı, komşu noktalar arasındaki korelasyonları hesaba katarak kovaryans matrisinin tahmininin sağlamlığını artırır ve böylece aykırı değerlerin veya aşırı güçlü örnekleme dalgalanmalarının etkisini sınırlamayı mümkün kılmaktadır. Böylece, son maksimum entropi kovaryans matrisi,

$$\hat{\Sigma}_{ME} = C + D \quad (16)$$

Şekildedir.

Bu yapı, matris C tarafından değiştirilen verilerin istatistiksel yapısını korurken, düzenleme D tarafından stabilize edilmiş iyi koşullandırılmış bir kovaryans matrisine sahip olmayı sağlamaktadır.

2.2.2 | Hibrit Kovaryans Tahmin (HCE) Yöntemi

Pamukçu ve çalışma arkadaşları (2015) tarafından geliştirilen Hibrit Kovaryans Tahmin Edicisi (Hybrid Covariance Estimator - HCE), maksimum entropinin avantajlarını, klasik stabilizasyon tekniklerini (Thomaz yöntemi dahil) ve verilerin yapısına uygun spesifik düzenlemeleri birleştirmektedir. Hesaplama süreci şu adımlardan oluşmaktadır.

- a) İlk tahmin: Maksimum olabilirlik veya maksimum entropi yöntemi ile hesaplanır.
- b) Thomaz stabilizasyonu: Kovaryans yapısını daha kararlı hâle getirmek için uygulanan işlemdir. Bu yapı aşağıdaki hesaplama adımlarından oluşturmaktadır. Burada, λ_i örneklem kovaryans matrisinin i'inci özdeğeri, $\bar{\lambda}$ özdeğerlerin ortalaması ve V ise özdeğerlere karşılık gelen özvektörlerdir.

$$\Lambda^* = \begin{bmatrix} \max(\lambda_1, \bar{\lambda}) & \cdots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \cdots & \max(\lambda_p, \bar{\lambda}) \end{bmatrix} \quad (17.a)$$

$$\hat{\Sigma}_{ME_STA} = V \Lambda^* V \quad (17.b)$$

maksimum entropinin sabitlenmesidir.

- c) Diğer kovaryans yapılarının dikkate alınması: Stabilize edilmiş matris, tahmin ediciyi rafine etmek için başlangıç noktası olarak kullanılır. Hibrit Kovaryans Tahmin Edicisi (HCE), önerilen kovaryans yapılarında (2.1) $\hat{\Sigma}$ yerine maksimum entropi yöntemi ile stabilize edilmiş şeklin (17.b) konulmasıyla elde edilmektedir.

Örneğin, kovaryans yapılarında (4.a), ... , (9.a) bu hibrit kovaryans tahmin edici biçimleri aşağıdaki Tablo 1’de özetlenmiştir. Burada, $\hat{\Sigma}_{ME_STA}$ modifiye edilmiş özdeğerler yöntemiyle stabilize edilmiş kovaryans matrisidir.

Kovaryans tahmin edicilerinin stabilizasyonu ve hibridizasyonunun kombinasyonu, özellikle küçük örneklem büyüklüğü veya yüksek boyut bağlamlarında daha güvenilir ve daha iyi koşullandırılmış kovaryans matrisleri elde etmeyi amaçlamaktadır. Birincisi, stabilizasyon, tahmin varyansını azaltmak, tekilliği önlemek ve daha iyi sayısal kararlılık sağlamak amacıyla ampirik kovaryans matrisinin spektrumunun özdeğerlerini ayarlayarak değiştirilmesinden oluşmaktadır. İkinci adımda, hibridizasyon bu stabilize matrisi kullanarak deneysel bilgi ve teorik yapıyı birleştiren analitik olarak yapılandırılmış tahmin ediciler (küresel, diyagonal, dışbükey, vb.) oluşurmaktadır. Bu ortak yaklaşım, ön yargı-varyans dengesinin kontrol edilmesine, kovaryans matrisinin tersine çevrilebilirliğinin sağlanmasına ve hibrit regresyon modelindeki (HRM) tahminlerin sağlamlığının ve performansının iyileştirilmesine olanak tanımaktadır.

Tablo 1: Hibrit Kovaryans Tahmin Edicileri (Stabilize Formlar)

Tahminci	Temel Formül	Hibrit Form ($\hat{\Sigma}, \hat{\Sigma}_{ME_STA}$ ile değiştirildi)
Empirical Bayes (EB)	$\hat{\Sigma}_{EB} = \hat{\Sigma} + \frac{(p-1)}{n.tr(\hat{\Sigma})} I_p$	$\hat{\Sigma}_{EB} = \hat{\Sigma}_{ME_STA} + \frac{(p-1)}{n.tr(\hat{\Sigma})} I_p$
Stipulated Ridge (SRE)	$\hat{\Sigma}_{SRE} = \hat{\Sigma} + p(p-1)[2n.tr(\hat{\Sigma})]^{-1} I_p$	$\hat{\Sigma}_{SRE} = \hat{\Sigma}_{ME_STA} + p(p-1)[2n.tr(\hat{\Sigma})]^{-1} I_p$
Stipulated Diagonal (SDE)	$\hat{\Sigma}_{SDE} = (1 - \pi)\hat{\Sigma} + \pi \text{diag}(\hat{\Sigma})$ $\pi = p(p-1)[2n.tr(\hat{\Sigma}^{-1}) - p]^{-1}$	$\hat{\Sigma}_{SDE} = (1 - \pi)\hat{\Sigma}_{ME_STA} + \pi \text{diag}(\hat{\Sigma})$ $\pi = p(p-1)[2n.tr(\hat{\Sigma}_{ME_STA}^{-1}) - p]^{-1}$
Convex Sum (CSE)	$\hat{\Sigma}_{CSE} = + \frac{n}{m+n} \hat{\Sigma} + (1 - \frac{n}{m+n}) [\frac{tr(\hat{\Sigma})}{p}] I_p$ $\beta = \frac{tr(\hat{\Sigma})^2}{tr(\hat{\Sigma}^2)}$	$\hat{\Sigma}_{CSE} = + \frac{n}{m+n} \hat{\Sigma}_{ME_STA} + (1 - \frac{n}{m+n}) [\frac{tr(\hat{\Sigma}_{ME_STA})}{p}] I_p$ $\beta = tr(\hat{\Sigma}_{ME_STA})^2 / tr((\hat{\Sigma}_{ME_STA})^2)$
Bozdogan Convex Sum (BCSE)	$\hat{\Sigma}_{BCSE} = \hat{p}\hat{\Sigma} + (1 - \hat{p})\hat{D}$ $\hat{p}=1/\alpha,$ $\alpha = (\frac{1}{n-1} \sum_{j=1}^p Var(x_j))$	$\hat{\Sigma}_{BCSE} = \hat{p}\hat{\Sigma}_{ME_STA} + (1 - \hat{p})\hat{D},$ $\hat{p}=1/\alpha_H,$ $\alpha = (1/(n-1)) \sum Var(x_j)$

2.3. ICOMP ve HCE’ye Dayalı Model Seçim Kriterleri ve Hibrit Regresyon Modeli (HRM)

2.3.1. Genel Giriş ve Bilgi Kriterleri

1970'lerden bu yana, en uygun modeli seçme sorunu istatistik ve modellemede merkezi bir yer tutmaktadır. Klasik seçim yaklaşımları esas olarak hipotez testlerine dayanırken, modern yöntemler uyum iyiliği ile modelin karmaşıklığı arasındaki dengeyi ölçmeyi amaçlayan bilgi kriterlerine dayanmaktadır. Bu kriterlerin temel ilkesi, yanıt değişkenlerinin gerçek dağılımı ile model tarafından tahmin edilen dağılım arasındaki Kullback–Leibler sapmasının (divergence) en aza indirilmesine dayanmaktadır. Akaike Bilgi Kriteri (AIC) model seçimi için ilk genel kriter olarak kabul edilen AIC, modelin maksimum olasılığı üzerinden uyum iyiliğini değerlendirirken, tahmin edilen parametre sayısını azaltmaktadır (Akaike, 1974). Formül aşağıdaki gibidir:

$$AIC = -2 \log L(\hat{\theta}) + 2k \quad (18)$$

“Burada $L(\hat{\theta})$ maksimum olabilirliği, k ise bağımsız parametre sayısını temsil etmektedir. Bu nedenle AIC, modelin verilere uyumu ile yapısal basitliği arasında bir denge kurmayı amaçlamaktadır. Aynı bağlamda, bu kriterin diğer varyantları da önerilmiştir; özellikle Schwartz 1978 tarafından geliştirilen Schwartz Bayesci Bilgi Kriteri (Schwarz's Bayesian Information Criterion-SBC) (Schwarz, 1978), ve Bozdoğan, 1987 tarafından geliştirilen Tutarlı yada Düzeltilmiş Akaike Bilgi Kriteri (CAIC) (Bozdoğan, 1987). Sırasıyla aşağıdaki gibi formüle edilmiştir:

$$SBC = -2 \log L(\hat{\theta}) + k \log(n) \quad (19)$$

$$CAIC = -2 \log L(\hat{\theta}) + k[\log(n) + 1] \quad (20)$$

Bu kriterler, karmaşık modellerin cezalarını güçlendirerek, özellikle örneklem büyüklüğü sınırlı olduğunda, tutumlu modellerin seçilmesini teşvik etmektedir. Ancak, etkili olmalarına rağmen, yukarıda belirtilen kriterlerin, açıklayıcı değişkenlerin yüksek korelasyonlu olduğu veya model yapısının doğrusal olmadığı durumlarda bazı sınırlamaları vardır. Bu gibi durumlarda, tahmin edicilerin kovaryans matrisi kararsız hâle gelebilir ve bu da Akaike veya Bayes kriterlerini modelin genel kalitesini değerlendirmek için yetersiz hâle getirebilmektedir. Bu eksikliklerin üstesinden gelmek için Bozdoğan, parametre kovaryansının karmaşıklığını açıkça içeren alternatif bir kriter olan Bilgi Karmaşıklığı Kriteri'ni (ICOMP) önermiştir (Bozdoğan, 1990). Bu kriter, ters Fisher bilgi matrisine (IFIM) dayanır ve yalnızca modelin uyum iyiliğini değil, aynı zamanda parametreler arasındaki bağımlılığın iç yapısını da değerlendirmeyi amaçlamaktadır. Genel biçimi şu şekildedir:

$$ICOMP = -2 \log L(\hat{\theta}) + 2C_1(\Sigma) \quad (21)$$

Burada Σ , tahmin edicilerin kovaryans matrisini gösterir ve $C_1(\Sigma)$, bu matrisin bilgi karmaşıklığını ölçmektedir. Bu karmaşıklık şu şekilde tanımlanmaktadır:

$$C_1(\Sigma) = \frac{p}{2} \log \left(\frac{\bar{\lambda}_a}{\bar{\lambda}_g} \right) \quad (22)$$

$\bar{\lambda}_a$ ve $\bar{\lambda}_g$, sırasıyla Σ 'nin özdeğerlerinin aritmetik ortalamasını ve geometrik ortalamasını temsil etmektedir. Bu iki ortalama önemli ölçüde farklı olduğunda, bu, özdeğerlerin yüksek bir dağılımına ve dolayısıyla modelin daha karmaşık olduğuna işaret etmektedir. Dolayısıyla, ICOMP hem uyum iyiliğini (olasılık yoluyla) hem de modelin iç yapısını (parametrelerin kovaryansı yoluyla) hesaba katarak AIC, BIC veya CAIC kriterlerinden daha eksiksiz bir görünüm sunmaktadır.

2.3.2 | Genişletme: HCE ve ICOMP'a dayalı hibrit regresyon modeli (HRM)

Pamukçu (2003), hibrit kovaryans tahmincilerine (HCE) dayalı hibrit regresyon modelini (HRM) geliştirerek bu yaklaşımları küçük örneklem büyüklüklerine ($n < p$) sahip durumlara genişletmiştir. Bu model, açıklayıcı değişkenlerin boyutunun gözlem sayısını aştığı regresyon modellerindeki tekillik ve sayısal istikrarsızlık sorunlarının üstesinden gelmektedir.

HRM Prosedür Adımları:

1. Tablo 1'de daha önce elde edilenler gibi farklı yapılaraya göre hibrit kovaryans matrisleri $\hat{\Sigma}_{HCE}$ 'nin tahmini.
2. Normal denklemler $X'X$ 'in çözümünde elde edilen matrislerin regresyon katsayılarını tahmin etmek için kullanılması.
3. Modellerin ICOMP kriter değerleri kullanılarak karşılaştırılması: Seçilen model, ICOMP'yi en aza indiren modeldir.

HRM modeli için klasik ve ICOMP kriterleri aşağıdaki formları almaktadır:

$$AIC(HRM) = n \log(2\pi) + n \log(\hat{\sigma}^2) + n + 2k \quad (23)$$

$$CAIC(HRM) = n \log(2\pi) + n \log(\hat{\sigma}^2) + n + k[\log(n) + 1] \quad (24)$$

$$ICOMP(HRM) = n \log(2\pi) + n \log(\hat{\sigma}^2) + n + 2C_1(Cov(\hat{\beta}_{HRM})) \quad (25)$$

Burada $\hat{\sigma}^2$ artıkların varyansını ve $Cov(\hat{\beta}_{HRM})$ potansiyel olarak yanlış belirlenmiş modelin katsayılarının tahmini kovaryans matrisini ifade etmektedir.

$$\hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2 \quad \text{'dır.}$$

$$S_k \text{ Çarpıklık katsayısı} = \frac{(\frac{1}{n} \sum_{i=1}^n \hat{\epsilon}_i^3)}{\hat{\sigma}^3} \quad \text{ve}$$

K_t Basıklık katsayısı = $\frac{(\frac{1}{n} \sum_{i=1}^n \hat{\varepsilon}_i^4)}{\hat{\sigma}^4}$ olmak üzere

$$Cov(\hat{\beta}_{HRM}) = \begin{bmatrix} \frac{\hat{\sigma}^2}{n} & 0 \\ 0 & \frac{2\hat{\sigma}^4}{n} \end{bmatrix} \begin{bmatrix} \frac{n}{\hat{\sigma}^2} & \frac{n S_k}{2 \hat{\sigma}^3} \\ \frac{n S_k}{2 \hat{\sigma}^3} & \frac{n(K_t-1)}{2\hat{\sigma}^4} \end{bmatrix} \begin{bmatrix} \frac{\hat{\sigma}^2}{n} & 0 \\ 0 & \frac{2\hat{\sigma}^4}{n} \end{bmatrix} \quad (26)$$

şeklinde tanımlanmaktadır.

3 | Simülasyon protokolü

Simülasyon protokolünün amacı, farklı kovaryans yapıları ve örneklem büyüklüklerine bağlı olarak hibrit regresyon modellerinin (HRM1–HRM5) karşılaştırmalı performanslarını değerlendirmektir. Bu değerlendirme, daha önce elde edilen ve Tablo 1’de sunulan hibrit modellere dayanmaktadır. Söz konusu beş model, sırasıyla HRM1’den HRM5’e karşılık gelmektedir. İki deneysel senaryo ele alınmıştır:

– Birinci senaryo, ($P_1 = 15$) bilgi verici (birbiriyle ilişkili) değişken ve ($P_2 = 25$) bilgi verici olmayan (bağımsız) değişken içermektedir; toplamda ($P = P_1 + P_2 = 40$) açıklayıcı değişken vardır.

– İkinci senaryo ise ($P_1 = 70$) bilgi verici ve ($P_2 = 30$) bilgi verici olmayan değişken içermektedir; dolayısıyla toplam ($P = 100$)’dür. Her bir senaryo için örneklem büyüklüğü (n), ilk durumda 5’ten 25’e, ikinci durumda ise 10’dan 50’ye kadar değişmektedir. Her parametre kombinasyonu, sonuçların istatistiksel kararlılığını sağlamak amacıyla γ kez tekrarlanmıştır. İlk aşamada $\gamma = 100$ yineleme yapılmış, ikinci aşamada ise $\gamma = 500$ ’e çıkarılmıştır. Böylece yineleme sayısındaki artışın, tahminlerin doğruluğu ve yakınsama üzerindeki etkisi incelenmiştir.

İlişkili değişkenler $X_1 \in \mathbb{R}^{n \times p_1}$, çok değişkenli normal dağılım $MVN(0, \Sigma)$ ’ye göre oluşturulmuştur; burada kovaryans matrisi Σ , bileşenler arasında $r=0,5$ ’te sabitlenmiş bir korelasyon sağlayacak şekilde oluşturulmuştur. Bağımsız değişkenler $X_2 \in \mathbb{R}^{n \times p_2}$, düzgün bir dağılım olan $U(0,1)$ ’i takip etmektedir. Hata vektörü ε , bir normal dağılımdan gelir ve hata varyansı $\sigma^2 = 0.25$ ’te sabitlenmiştir. Simüle edilen model aşağıdaki gibi yazılmıştır:

$$y = X\beta + \varepsilon\sigma$$

Burada $X=[X_1, X_2]$ açıklayıcı değişkenlerin tam matrisini ve β katsayı vektörünü ifade etmektedir. Simüle edilmiş verilerden, ampirik kovaryans matrisleri ($\hat{\Sigma}_{MLE}$) ve büzülme tipi

tahmin ediciler ($\hat{\Sigma}_{ME}$) hesaplanmaktadır. İkincisi, örnek matrisinden ve diyagonal hedef matrisinden gelen bilgileri, büzülme yoğunluğu λ adı verilen optimal bir ağırlıklandırma faktörü ile birleştiren Ledoit ve Wolf (2004) prosedürünü kullanmaktadır. Simülasyonlar sırasında, korelasyon matrisi ve varyans vektörü için tahmin edilen optimal yoğunluk değerleri sırasıyla $\lambda \approx 0.82-0.87$ ve $\lambda.var \approx 0.05-0.09$ aralığında bulunmuştur. Bu değerler, orta düzeyde korelasyona sahip kovaryans yapılarıyla uyumlu, ölçülü bir düzenlileştirme derecesine işaret etmekte olup, sayısal kararlılığı ve tahmin edicilerin pozitif tanımlılığını sağlamak amacıyla tahmin sürecinin her yinelenmesinde otomatik olarak belirlenmektedir. Böylece söz konusu parametreler, ampirik kovaryans matrisi ile düzenlileştirme aşamasında kullanılan hedef matris arasındaki uyum derecesini yansıtan bir rol üstlenmektedir.

Önceden tanımlanan kovaryans yapıları, beş hibrit modelin (HRM1–HRM5) her birine entegre edilmiştir. Bu modellerin her biri için, regresyon katsayıları, artık varyansı ve bilgi ölçütleri (AIC, CAIC ve ICOMP) hesaplanmış, ardından daha kararlı tahminler elde etmek amacıyla tüm yinelenmeler (tekrarlamalar) üzerinden ortalaması alınmıştır. Her bir örneklem büyüklüğü (n) ve senaryo kombinasyonuna karşılık gelen toplu sonuçlar, aşağıdaki tablolarda sunulmaktadır.

Tablo 2: Senaryo 1 : P1 = 15, P2 = 25

n	HRM	AIC	CAIC	ICOMP
5	HRM1	115.97454	140.3521	4997.27094
5	HRM2	112.48808	136.8656	1234.02990
5	HRM3	129.97501	154.3525	72894.23547
5	HRM4	93.86045	118.2380	55.85124
5	HRM5	115.97454	140.3521	4997.27094
15	HRM1	197.0522	265.3742	2053.51294
15	HRM2	176.8792	245.2012	421.62661
15	HRM3	222.7950	291.1170	9340.02402
15	HRM4	100.7961	169.1181	23.42648
15	HRM5	197.0522	265.3742	2053.51294
20	HRM1	230.3147	310.1440	1251.0353
20	HRM2	202.0838	281.9131	303.2693
20	HRM3	262.3725	342.2018	5077.3317
20	HRM4	103.2302	183.0595	24.7965
20	HRM5	230.3147	310.1440	1251.0353
25	HRM1	268.5342	357.2892	996.01737
25	HRM2	232.5034	321.2585	289.39324
25	HRM3	309.9892	398.7443	3923.32863
25	HRM4	107.9551	196.7102	29.15669

25	HRM5	268.5342	357.2892	996.01737
----	------	----------	----------	-----------

Tablo 3: Senaryo 2 : $P1 = 70, P2 = 30$

n	HRM	AIC	CAIC	ICOMP
10	HRM1	304.4015	434.6600	60883.1536
10	HRM2	290.8433	421.1018	9465.4170
10	HRM3	326.0947	456.3532	509344.2939
10	HRM4	243.4567	373.7152	194.2874
10	HRM5	304.4133	434.6718	61032.9608
20	HRM1	403.2596	602.8329	20936.8864
20	HRM2	375.4133	574.9865	3649.0802
20	HRM3	434.1411	633.7143	96713.6529
20	HRM4	282.3817	481.9549	134.6366
20	HRM5	403.2754	602.8486	20959.1510
25	HRM1	447.1053	668.9929	12368.2053
25	HRM2	414.7786	636.6662	2530.8759
25	HRM3	484.8601	706.7476	54276.1158
25	HRM4	300.1439	522.0315	134.0887
25	HRM5	447.1109	668.9985	12372.1280
50	HRM1	655.9047	947.1070	3115.4169
50	HRM2	607.8384	899.0407	1181.7136
50	HRM3	738.7845	1029.9868	14051.2592
50	HRM4	385.8166	677.0189	197.3385
50	HRM5	655.9047	947.1070	3115.4169

Tablo 4: Senaryo 2'nin Sonuçları: $\gamma = 500$ olduğunda $P1 = 70, P2 = 30$

n=10

HRM	AIC	CAIC	ICOMP
HRM1	304.3076	434.5661	59470.0667
HRM2	290.5602	420.8187	9232.3386
HRM3	325.4797	455.7382	493307.8548
HRM4	242.1647	372.4232	177.6246
HRM5	304.3289	434.5874	59591.6895

n=20

HRM	AIC	CAIC	ICOMP
HRM1	403.4696	603.0428	20681.9701

HRM2	376.1507	575.724	3704.0345
HRM3	434.8356	634.4089	97912.6927
HRM4	285.5137	485.0869	142.6733
HRM5	403.4859	603.0591	20695.1161

n=25

HRM	AIC	CAIC	ICOMP
HRM1	447.4935	669.3811	12605.958
HRM2	414.9249	636.8125	2555.7025
HRM3	484.8661	706.7536	54901.0567
HRM4	301.0545	522.9421	136.6469
HRM5	447.5016	669.3892	12609.4036

n=50

HRM	AIC	CAIC	ICOMP
HRM1	653.8291	945.0314	3026.9123
HRM2	606.2232	897.4255	1163.2513
HRM3	737.5585	1028.7608	13960.8596
HRM4	384.0572	675.2595	195.2684
HRM5	653.8379	945.0402	3027.2993

4 | Sonuçların Yorumlanması ve Tartışma

İki simülasyon senaryosundan — Senaryo 1 ($P_1 = 15$, $P_2 = 25$) ve Senaryo 2 ($P_1 = 70$, $P_2 = 30$) elde edilen sonuçlar, farklı örneklem büyüklükleri (n) ve hibrit kovaryans yapıları altında beş farklı Hibrit Regresyon Modeli (HRM) varyantının karşılaştırmalı davranışını açıkça ortaya koymaktadır. Senaryo 1’de, toplam açıklayıcı değişken sayısının ($P = 40$) orta düzeyde kaldığı durumda, HRM4 ve HRM2 modelleri sistematik olarak AIC, CAIC ve ICOMP kriterleri bakımından en düşük değerlere ulaşmıştır. Bu düşük değerler, modellerin veriye daha iyi uyum sağladığını ve aşırı parametreleşme riskini sınırladığını göstermektedir.

Özellikle, MLE–stabilize–BCSE kombinasyonuna dayanan HRM4 modeli, tüm senaryolar ve bilgi kriterleri için en kararlı ve en düşük skorları elde etmiştir. Bu durum, stabilize edilmiş bir kovaryans matrisine uygulanan BCSE düzenleme yönteminin (regularizasyonunun), küçük örneklem büyüklüğüyle karakterize edilen durumlarda yanlılık–varyans dengesini (bias–variance trade-off) etkili bir biçimde optimize ettiğini göstermektedir. Senaryo 3’te, yüksek boyutlulukla ($P = 100$) ve düşük ila orta düzeyde örneklem büyüklüğüyle ($n \leq 50$) karakterize edilen durumda, yinleme sayısı (γ) artsa bile modellerin genel performans

düzeninin değişmeden kaldığı gözlemlenmiştir. Bununla birlikte, yineleme sayısındaki artış, tüm bilgi kriterleri için daha düşük değerlere karşılık gelen daha hassas tahminlerin elde edilmesini sağlamaktadır. HRM2 ve HRM4 modelleri, diğer varyantlara kıyasla üstün performanslarını koruyarak yüksek boyutlu ortamlara uyum sağlama kapasitelerini doğrulamaktadır. Buna karşılık, sırasıyla maksimum entropi kovaryansına (ME) ve onun düzeltilmemiş versiyonlarına dayanan HRM3 ve HRM1 modelleri, belirgin biçimde daha yüksek ICOMP değerleri sergilemektedir. Bu durum, boyut arttıkça kovaryans parametrelerinin tahmininde ortaya çıkan istikrarsızlıktan kaynaklanan artmış bilgi karmaşıklığını yansıtmaktadır.

HRM4 modelinde dengeleme (stabilizasyon) ve hibritleştirme tekniklerinin birlikte uygulanması, tahmin kalitesinin iyileştirilmesinde özellikle belirleyici bir rol oynamaktadır. HCE (Hybrid Covariance Estimator) türü yöntemlerle kovaryans matrisinin stabilize edilmesi, küçük örneklem büyüklükleri veya kötü koşullanmış tasarım matrisleriyle ilişkili kararsızlık (instabilite) sorunlarını düzeltmeyi mümkün kılmaktadır. Buna paralel olarak, regresyon modelinin hibritleştirilmesi, alt modellerden (HRM1–HRM3) türetilen bileşenlerin katkısını ağırlıklandırarak düzenleyici (regularizasyon) bir mekanizma gibi işlev görür; böylece tahmin edicilerin aşırı varyansı sınırlandırılmaktadır. Bu birleşim, yanlılık–varyans (bias–variance) arasında optimal bir denge sağlayarak parametre tahminlerinde daha yüksek dayanıklılık (robustluk) ve istikrar (stabilite) elde edilmesini mümkün kılmaktadır. Ayrıca, Pamukçu 2015’te tarafından yürütülen çalışma sonucunda, en uygun kovaryans yapısının p/n oranına sıkı biçimde bağlı olduğu ortaya konmuştur. Daha açık bir ifadeyle, $p/n \leq 5$ durumunda MLE/STA/CSE kombinasyonu en iyi performansı gösterirken, $p/n > 5$ olduğunda MLE/STA/BCSE yapısının daha başarılı sonuçlar verdiği belirlenmiştir (Pamukçu, 2015). Elde edilen sonuçlar ayrıca bilgi kriterlerinin, özellikle de ICOMP kriterinin, modellerin yapısal stabilitesinin değerlendirilmesinde belirleyici bir rol oynadığını doğrulamaktadır. Kovaryans matrisinin bilgi karmaşıklığını açık biçimde dikkate alan ICOMP, kötü koşullanmış modelleri daha net biçimde ayırt etme olanağı sunarak model seçim sürecinin güvenilirliğini artırmaktadır.

5 | Sonuç ve Gelecek Perspektifleri

Elde edilen sonuçlar, spektral dengeleme (stabilizasyon) ve hibritleştirmeyi birleştiren HRM4 modelinin, kovaryans matrisinin kararsız veya kötü koşullanmış olduğu durumlarda doğrusal regresyon tahmini için sağlam (robust) ve yüksek performanslı bir yaklaşım sunduğunu doğrulamaktadır. Bu strateji, hem sayısal kararlılığı (nümerik stabiliteyi) hem de uyum kalitesini artırmakta, aynı zamanda bilgi karmaşıklığını azaltmaktadır.

Gelecekteki çalışmaların amacı, analizi çok yüksek boyutluluk ($p \gg n$) bağlamlarına genişletmek, HRM modelini diğer düzenlenleştirilmiş (regularize edilmiş) tahmin edicilerle —örneğin Ledoit–Wolf ve Graphical Lasso karşılaştırmak ve modelin ceza (penalizasyon) parametrelerine duyarlılığını incelemek olacaktır. Ayrıca, gerçek veri kümeleri üzerinde yapılacak doğrulama, simülasyonlarda gözlemlenen sağlamlığı teyit etmeye ve modelin genelleme kapasitesini değerlendirmeye olanak tanıyacaktır.

KAYNAKLAR

- Akaike, H. (1974). A new look at the statistical model identification. *IEEE Transactions on Automatic Control*, 19(6), 716–723. <https://doi.org/10.1109/TAC.1974.1100705>
- Bozdogan, H. (1987). ICOMP: A new model-selection criterion. In *Proceedings of the 1st Conference of the International Federation of Classification Societies* (pp. 599–608).
- Bozdogan, H. (1990). On the information-based measure of covariance complexity and its application to the evaluation of multivariate linear models. *Communications in Statistics – Theory and Methods*, 19(1), 221–278. <https://doi.org/10.1080/03610929008830199>
- Bozdogan, H. (2009). A new class of information complexity (ICOMP) criteria with an application to customer profiling and segmentation. *İstanbul Üniversitesi İşletme Fakültesi Dergisi*, 39(2), 370–398. <https://dergipark.org.tr/en/pub/iuisletme/issue/9248/115706>
- Donoho, D. L. (2000). *High-dimensional data analysis: The curses and blessings of dimensionality*. Stanford University.
- Fiebig, D. G. (1984). On the maximum-entropy approach to undersized samples. *Applied Mathematics and Computation*, 14(3), 301–312. [https://doi.org/10.1016/0096-3003\(84\)90027-4](https://doi.org/10.1016/0096-3003(84)90027-4)
- Haff, L. R. (1980). Empirical Bayes estimation of the multivariate normal covariance matrix. *The Annals of Statistics*, 8(3), 586–597. <https://www.jstor.org/stable/2240594>
- Hoerl, A. E., & Kennard, R. W. (1970). Ridge regression: Biased estimation for nonorthogonal problems. *Technometrics*, 12(1), 55–67. <https://doi.org/10.1080/00401706.1970.10488634>
- Hoerl, R. W. (2020). Ridge regression: A historical context. *Technometrics*, 62(4), 420–425.
- Jaynes, E. T. (1957). Information theory and statistical mechanics. *Physical Review*, 106(4), 620–630. <https://doi.org/10.1103/PhysRev.106.620>

- Ledoit, O., & Wolf, M. (2004). A well-conditioned estimator for large-dimensional covariance matrices. *Journal of Multivariate Analysis*, 88(2), 365–411. [https://doi.org/10.1016/S0047-259X\(03\)00096-4](https://doi.org/10.1016/S0047-259X(03)00096-4)
- Pamukçu, E. (2015). Aşırı derecede küçük örneklem problemi için hibrit regresyon modeli. *Celal Bayar University Journal of Science*, 13(3), 803–813.
- Pamukçu, E., Bozdoğan, H., & Çalık, S. (2015). A novel hybrid dimension reduction technique for undersized high dimensional gene expression data sets using information complexity criterion for cancer classification. *Computational and Mathematical Methods in Medicine*, 2015(1), 370640. <https://doi.org/10.1155/2015/370640>
- Press, S. J. (1975). *Estimation of a normal covariance matrix*. Rand Corporation.
- Schwarz, G. (1978). Estimating the dimension of a model. *The Annals of Statistics*, 6(2), 461–464. <https://www.jstor.org/stable/2958889>
- Shurygin, A. (1983). The linear combination of the simplest discriminator and Fisher's one. *Applied Statistics*, 32, 144–158.
- Stein, C. (1975). Estimation of a covariance matrix. In *Proceedings of the 39th Annual Meeting of the Institute of Mathematical Statistics*, Atlanta, GA, USA. <https://cir.nii.ac.jp/crid/1571417125492037632>
- Tibshirani, R. (1996). Regression shrinkage and selection via the lasso. *Journal of the Royal Statistical Society: Series B (Methodological)*, 58(1), 267–288. <https://doi.org/10.1111/j.2517-6161.1996.tb02080.x>
- Yilun Chen, Wiesel, A., & Hero, A. O. (2011). Robust shrinkage estimation of high-dimensional covariance matrices. *IEEE Transactions on Signal Processing*, 59(9), 4097–4107. <https://doi.org/10.1109/TSP.2011.2138698>
- Zou, H., & Hastie, T. (2005). Addendum: Regularization and variable selection via the elastic net. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 67(5), 768. <https://doi.org/10.1111/j.1467-9868.2005.00527.x>