

Predicting Age based on Bio-measurements

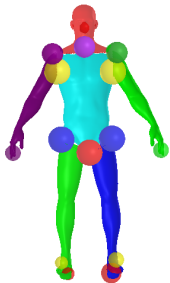
Peyton Ardoin, Khalil El-abbassi, Aditya Baisakh, Marilin Guerrero Laos, Tre' Henson, Matthew Lemoine, Austin Louque, Efren Mesino Espinosa, Aduragbemi Oyerinde, Shalini Shalini, Gowri Priya Sunkara

August 18, 2026

- 1 History and Motivation
- 2 Data Cleaning
- 3 Models on Penn Data
- 4 Models on CAESAR Dataset
- 5 Goals and Plans Moving Forward

History and Motivation

History of the Project



In previous projects, the Math Consultation Clinic has been interested in predicting specific anthropometric values based on certain biomarkers. The DXA scan was the standard method of obtaining these values, but previous projects were able to predict these values with great precision using machine learning algorithms.

Figure 1: An example of the avatar taken of the patient. This is where we get measurements of the body.

History of the Project

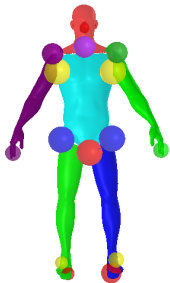


Figure 1: An example of the avatar taken of the patient. This is where we get measurements of the body.

In previous projects, the Math Consultation Clinic has been interested in predicting specific anthropometric values based on certain biomarkers. The DXA scan was the standard method of obtaining these values, but previous projects were able to predict these values with great precision using machine learning algorithms.

In these other projects, a scan of a patient's body is taken. Then from this scan measurements can be obtained and used in the training and testing of a machine learning algorithm.

This current project to predict a patient's biological age is motivated by the success of previous projects. A patient's biological age can provide a great deal of insight into their overall health. If, for example, you have a 23 year old with the biological age of a 30 year old, we would expect that this patient has some health complications that need to be addressed.

This current project to predict a patient's biological age is motivated by the success of previous projects. A patient's biological age can provide a great deal of insight into their overall health. If, for example, you have a 23 year old with the biological age of a 30 year old, we would expect that this patient has some health complications that need to be addressed.

The motivation behind this project is to use different measurements from a given patient's body to predict their age. This prediction made is the *biological age* of the patient.

Data Cleaning

Originally, the dataset provided contained over 120 biomarkers and data points for each patient. The biomarkers contain information about both individual body parts or the body as a whole (Ex. Torso Volume vs Total Volume).

Originally, the dataset provided contained over 120 biomarkers and data points for each patient. The biomarkers contain information about both individual body parts or the body as a whole (Ex. Torso Volume vs Total Volume).

Any other data points that were not required anthropometric biomarkers were removed. After removing those, as well as any redundant columns, the dataset contains 94 usable biomarker columns for each patient.

Issues in model training can occur if there are any unintentional empty fields in the data. There were 227 instances of incomplete patient data, which means that these patients could not reliably be used for training, and thus were removed from the dataset. This left 2164 rows of patient data that we could use for our training and testing.

Further Cleaning

Issues in model training can occur if there are any unintentional empty fields in the data. There were 227 instances of incomplete patient data, which means that these patients could not reliably be used for training, and thus were removed from the dataset. This left 2164 rows of patient data that we could use for our training and testing.

The exception to this removal is a column titled "-Bust/Chest Circumference Under Bust (mm)". Around half of our patients were missing data in this column, far too many to be removed. After further inspection, this was due to this data not being collected for males, so removal was not necessary.

Due to this biomarker being missing from all male patients, the dataset was split into 6 different datasets. These datasets were:

- Both Genders with/without Bust Measurement
- Males with/without Bust Measurement
- Females with/without Bust Measurements

Due to this biomarker being missing from all male patients, the dataset was split into 6 different datasets. These datasets were:

- Both Genders with/without Bust Measurement
- Males with/without Bust Measurement
- Females with/without Bust Measurements

Not only does this ensure that model training can correctly run without any missing data issues, but it allows us to further probe the data for any gender specific patterns that may influence a model's decision making.

Models on Penn Data

Models that we are using on Penn Data

The following are the models that we have used:

- Linear Regression (LR), Ridge Regression (RR) and Bayesian Ridge Regression (BRR)
- Random Forest (RF)
- XGBoost

Linear Regression (PENN)

Sub-Models Studied

Linear Regression (LR)

LR strives to find an optimal linear relationship between several input biomarkers and chronological age of a patient.

Ridge Regression (RR)

Ridge Regression improves this by shrinking the coefficients to reduce overfitting, especially when biomarkers are correlated.

When biomarkers are highly correlated (e.g., waist and narrow waist), Ridge Regression shrinks their coefficients to reduce overfitting.

Bayesian Ridge Regression (BRR)

Bayesian Ridge Regression goes one step further by treating the coefficients as probability distributions, allowing it to account for uncertainty in the data and often produce more stable predictions.

Linear Regression (PENN)

Feature Selection for Body Shape Biological Age (BSBA)

46 Biomarkers



Feature Selection

Ridge coefficient-based
reduction



21 Biomarkers

**Reduced Biomarkers
with Improved
RMSE**

Final BSBA Biomarker Groups

Morphometric Biomarkers

- Waist, narrow waist
- Thigh, calf, and leg measurements
- Inseam and outside leg length
- Arm length and arm volume
- Leg surface area

Body Composition Biomarkers

- Appendicular Lean Mass (ALM)
- Bone Mineral Density (BMD)
- Total Body Volume
- Torso Volume

Comparison with Previous BSBA Models

Model	Biomarkers	Overall RMSE
Bayesian Ridge	46	13.278
Ridge	46	13.320
Bayesian Ridge	21	12.812
Ridge	21	12.867

Conclusion

Reducing the model from **46 biomarkers** to **21 biomarkers** improved overall RMSE and produced a simpler, more interpretable BSBA model.

Best overall model: Bayesian Ridge with 21 biomarkers

- Random Forest is built on the base learner of decision trees
- At each node/tree a single marker and a single threshold is chosen. This threshold is chosen to reduce the impurities of the target. This is recurring at each tree until the maximum depth is reached.
- The main parameters for random forest is the number of estimators used, the max depth of tree, and the number of samples used per decision tree

- Use a random search algorithm to test combinations of the many parameters used in RF across a grid of different values selected by the developer
- Train the RF model on 80% of the data to give it enough information to make reliable decisions.
- Test the training on a tiered list of markers of high, medium, and low importance split into thirds based on value of any given marker.
- Also experimented with training based on interacting markers.

Random Forest Results

Below are the results obtained from testing Random Forest.

Model	Error
Random Forest with 11 age brackets	11.6 years
Random Forest Tuned Hyper-parameters	10.7 years
Random Forest Bracketed markers	13.3 years
Random Forest Bracketed Markers with interactions	12.42 years

Table 1: Results from RF tests.

Random Forest Results

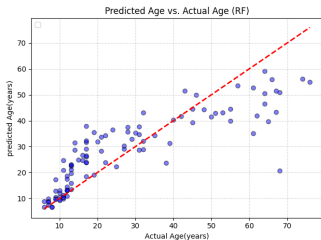


Figure 2: RF with tuned parameters and 11 age brackets

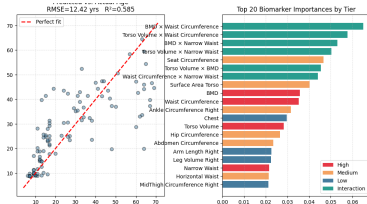


Figure 3: RF with tuned parameters and interacting tiered markers

Random Forest Limitations

- RF already assigns markers weights and importance values as it trains each node so inputting rigid values for the importance of a marker limits the models ability to treat them as continuous variables and make a smooth model.
- In my testing of RF all RMSE values were relatively high(above 10 years).
- RF consistently over estimates ages withing the 10 to 40 range and underestimates ages 45+. This leaves many questions about the interpretation of aging in regards to the physical measurements we train the models on.

e**X**treme **G**radient **B**oosting - also an ensemble algorithm but uses boosting vs. bagging (sequential vs. parallel)

RMSE 12.88

Root mean squared error, in years.

Observations

- Model performed best on the densely-sampled **5 - 12 age group** (RMSE ≈ 7.3)
- Error concentrates at the sparse **extremes**: teens and, especially, older adults (only 13 patients, RMSE > 24)
- These thinly-sampled age bands inflate the overall RMSE disproportionately

Models on CAESAR Dataset

The following are the models that we have used:

- Multiple Linear Regression (MLR) with Principal Component Analysis (PCA)
- Support Vector Regression (SVR)
- Least Squares Support Vector Regression (LSSVR)
- Extreme Gradient Boosting (XGBoost)
- Comparison based on Groups with Kernel-based regression models

Linear Regression (CAESAR)

Feature Selection and Extraction on the CAESAR Dataset

Motivation

- Previous models on the Penn dataset employed Ridge Regression (RR) and Bayesian Ridge Regression (BRR) for feature selection for Body Shape Biological Age (BSBA).
- This improving prediction accuracy (minimizing RMSE) while simplifying the model.
- However, a global regression model may not capture age-dependent relationships between biomarkers, so the below methods were implemented on the CAESAR dataset.

Linear Regression (CAESAR)

Methodology

Objective

Predict chronological age from anthropometric biomarkers using segmented Multiple Linear Regression (MLR) via Correlation Analysis and Principal Component Analysis (PCA).

Model Pipeline

- 1 Use the cleaned anthropometric biomarker dataset.
- 2 Perform correlation and variance inflation factor (VIF) analysis to examine and target multicollinearity.
- 3 Standardize biomarker measurements.
- 4 Apply PCA within each age group, retaining approximately 95% of the total variance.
- 5 Train an independent MLR model for each age group (as shown in the results).
- 6 Evaluate model performance using MAE, RMSE, and R^2 .

Linear Regression (CAESAR)

Two Methods to Enhance Model Performance

Definition

Correlation analysis is a statistical method used to infer the strength and direction of the relationship between two or more variables. Key concepts related to this analysis in regression models are **multicollinearity**—when independent variables are highly mutually correlated—and **variance inflation factor (VIF)**—a metric used to detect *by how much* that correlation inflates the variance of coefficients in the regression.

Definition

Principle component analysis (PCA) is a dimensionality reduction technique in which the dataset is first divided into predefined age groups, and PCA is performed independently within each group.

Linear Regression (CAESAR)

Results of Correlation and VIF Analysis

Below shows the top highly correlated feature/biomarker pairs.

TOP HIGHLY CORRELATED FEATURE PAIRS		
Infraorbitale Ht Lt Stand (mm)	Infraorbitale Ht Rt Stand (mm)	0.999
Infraorbitale Ht Sit Lt (mm)	Infraorbitale Ht Sit Rt (mm)	0.997
Thumb Tip Reach (mm)	TTR 2 (mm)	0.995
Weight (kg)	Volume	0.995
Thumb Tip Reach (mm)	TTR 3 (mm)	0.994
Thumb Tip Reach (mm)	TTR 1 (mm)	0.992
Knee Ht Stand Lt (mm)	Knee Ht Stand Rt (mm)	0.991
Leg Volume Left	Leg Volume Right	0.991
Acromial Ht Stand Lt (mm)	Acromial Ht Stand Rt (mm)	0.990
Axilla Ht Lt (mm)	Axilla Ht Rt (mm)	0.989
Cervicale Ht (mm)	Suprasternale Ht (mm)	0.989
Infraorbitale Ht Rt Stand (mm)	Suprasternale Ht (mm)	0.988
Infraorbitale Ht Lt Stand (mm)	Suprasternale Ht (mm)	0.988
Acromial Ht Stand Lt (mm)	Cervicale Ht (mm)	0.987
Infraorbitale Ht Lt Stand (mm)	Stature (mm)	0.987
Infraorbitale Ht Rt Stand (mm)	Stature (mm)	0.987
Cervicale Ht (mm)	Stature (mm)	0.987
Acromial Ht Stand Rt (mm)	Cervicale Ht (mm)	0.986
TTR 2 (mm)	TTR 3 (mm)	0.986
Torso Volume	Volume	0.986

The top *finite-valued* VIF-valued features are as listed:

Volume, Surface Area Total, Torso Volume, Surface Area Torso, Surface Area Leg Right/Left, Leg Volume Right/Left, Surface Area Arm Left/Right, Arm Volume Left/Right, Thumb Tip Reach, and Infraorbitale Ht Rt/Lt Stand.

Linear Regression (CAESAR)

How PCA Improves the Regression Model

Motivation

Correlation and VIF analysis showed that many biomarkers measured similar underlying anatomical characteristics, leading to multicollinearity.

Principal Component Analysis (PCA)

- Standardize all biomarker measurements.
- Compute principal components as orthogonal linear combinations of the biomarkers.
- Retain only the components explaining approximately 95% of the total variance.
- Train a separate MLR model using the PCA components within each age group.

Procedure: **Standardize Biomarkers** → **PCA**

→ ***linearly independent*** Principal Components (PCs) →
Segmented MLR Model

Linear Regression (CAESAR)

MLR w/ PCA Results

Age Group	Samples	No. of PCs	Variance	RMSE
0–20	58	17	95.36%	0.56
20–30	539	27	95.20%	2.69
30–40	612	26	95.01%	2.88
40–50	639	27	95.11%	2.87
50–60	396	28	95.05%	2.86
60–70	142	25	95.29%	1.85

Overall Performance and Findings

MAE = 2.331 years
RMSE = 2.746 years
MSE = 7.538 years²
R² = 0.9486

- PCA reduced the dimensionality of the biomarker space while retaining approximately 95% of the variance within each age group.
- This pipeline achieved an overall RMSE of 2.75 years and an R^2 of 0.9486, substantially improving upon the initial performance.

Linear Regression (CAESAR)

MLR w/ PCA Results (*contd.*)

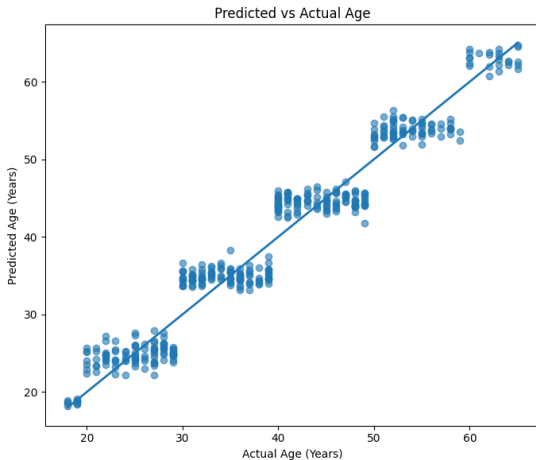


Figure 4: Segmented MLR w/ PCA Model; $RMSE = 2.75$ years

XGBoost (CAESAR)

Improvements over the original 12.88 RMSE global prediction error from first dataset due to:

- Caesar dataset having 4 times as many patients
- Splitting sexes removes some physical discrepancy noise that was unrelated to age-based changes
- Hyperparameters tuned via 5-fold CV

Model	Error
XGBoost w/ tuned hyperparameters (Males)	8.24 years
XGBoost w/ tuned hyperparameters (Females)	8.66 years

Table 2: Results from XGBoost tests.

Support Vector Regression

- Support vector regression is a regression model that predicts continuous values by fitting the simplest function to the data
- SVR is built on an ϵ -tube. This tube is created by fitting two planes to the data. The distance, ϵ , between the edges of the ϵ tube is a tunable parameter.
- SVR uses the point on the edges and the outside of the ϵ -tube to create the model. These points are known as support vectors.
- These support vectors train the model to make the simplest function to fit the data.
- Each of these points is assigned a weight depending on the parameter C . This controls how much each support vector can pull the model depending on its distance from the ϵ -tube

- Use a single global model to get the true predictions
- Determine important and unimportant biomarkers.
- Drop unimportant markers from final model based on SHAP value.
- Optimize parameters through random search algorithm to test combinations of C and ϵ .
- Pair the SVR with K-Means clustering to compare clusters of representatives to other patients of similar age groups
- Fight regression dilution through a corrected, uncompressed prediction based on the training set to get an unbiased per patient prediction. Centered at zero instead of the model mean.
- Run a centile analysis on the data to determine the atypicality of each test patient and their markers to see why the model is predicting them higher or lower than their chronological age.

Using Support vector regression we have collected a number of different results through various approaches using this model. Below is a table highlighting some of the results.

Input	Result
Whole data set split into 10 year age columns	7.9 years
SHAP tuned model (0.02)	2.4 years
SHAP tuned model with tuned parameters	1.94 years
SHAP tuned model with tuned parameters and tuned SHAP thresholds	1.92 years
New data set with SHAP and K-Means on one age bracket	8.54 years height

SVR Results

Highlighted below are some of the plots from the recent SVR runs

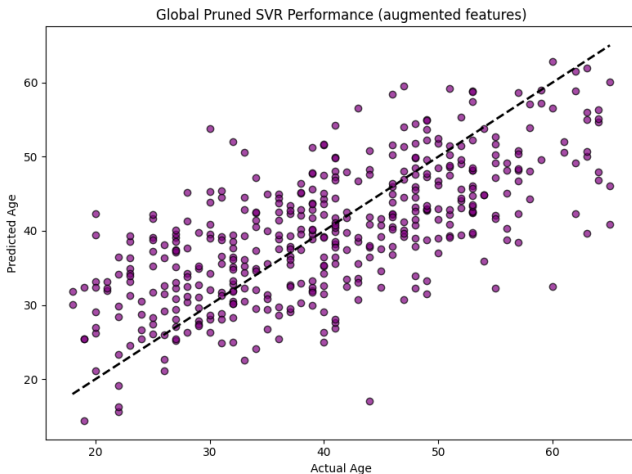


Figure 5: SVR SHAP with K-Means; RMSE=8.54 years

Below is a snip-it of the output data table

Table 3: Results from SVR run output to CSV

PPT ID	Actual Age	Corrected Gap	Centile Deviation
478	42 y/o	-1.5 years	-13.9
1745	51 y/o	-11.3 years	-24.8
1851	64 y/o	-9.1 years	-11.9

SVR Results

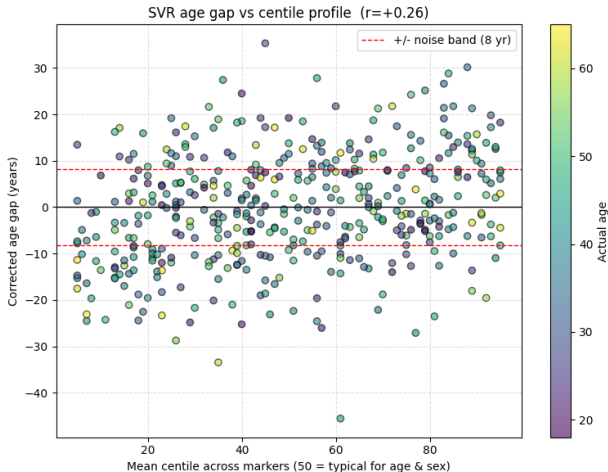


Figure 6: Age gap Vs. centile deviation

SVR Results

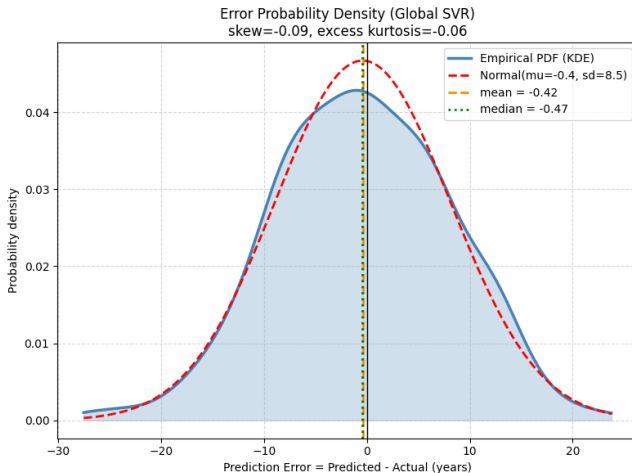


Figure 7: Probability density of Age Gap

- Interpreting outputs
 - in the pursuit of lower RMSE it continues to be difficult to interpret the outputs on a per patient status. Finding what error is model noise and what is a true biological indicator
- Dataset Limitations
 - In current literature, most Bio-age projects have an output variable in their dataset. Something like mortality or morbidity to fit a rate of true aging to the patients.

Least Squares Support Vector Regression

Least Squares Support Vector Regression (LSSVR) is similar to Support Vector Regression in that they have the same ϵ tube around a hyperplane in your space and C parameter. In SVR, we determine the placement of the tube by looking at the distances of the support vectors outside of the tube and minimizing. In LSSVR, we look at the least squared distance of the support vectors outside the tube to minimize.

Results and Methods for LSSVR

Using the LSSVR model we have obtained varied results. Below are average results using the 5-fold cross validation with the LSSVR method.

Input	Results
Caesar (Female)	8.37 year
Caesar (Male)	8.06
Caesar (whole)	8.20
Caesar (whole; no gender)	8.21
Whole Dataset (Pennington)	11.73

Table 4: Initial attempts and results for the LSSVR with the Caesar dataset compared to the Pennington Dataset.

Kernel-Based Methods

Kernel-based methods use similarity functions called kernels to assign labels to unlabeled data by comparing them to previous training data instead of a fixed set of parameters. This "kernel trick" allows the algorithm to operate in high-dimensional spaces by looking at the inner product of the input vectors without calculating explicit coordinates in those spaces, which is computationally cheaper.

- 1 **Kernel Ridge Regression:** Ridge regression combined with a kernel to model linear and nonlinear relationships.
- 2 **Support Vector Regression:** Regression of data by learning nonlinear relationships in a high-dimensional feature space to predict continuous target values.
- 3 **Gaussian Process Regression:** Probability distribution of possible functions fitting the data, using a kernel to build on previous tests.

RBF Regression Models

Shared Radial Basis Function (RBF) Kernel

$$K(\mathbf{x}_i, \mathbf{x}_j) = \exp\left(-\gamma \sum_{m=1}^p w_m (x_{im} - x_{jm})^2\right),$$

where w_m is the importance assigned to biomarker m .

Support Vector Regression (SVR)

$$f(\mathbf{x}) = \sum_{i=1}^N (\alpha_i - \alpha_i^*) K(\mathbf{x}_i, \mathbf{x}) + b$$

Kernel Ridge Regression (KRR)

$$\hat{y}(\mathbf{x}) = \sum_{i=1}^N \alpha_i K(\mathbf{x}_i, \mathbf{x})$$

$$\alpha = (K + \lambda I)^{-1} \mathbf{y}$$

Gaussian Process Regression (GPR)

$$f(\mathbf{x}) \sim GP(m(\mathbf{x}), k(\mathbf{x}, \mathbf{x}'))$$

$$k(\mathbf{x}, \mathbf{x}') = \sigma_f^2 K(\mathbf{x}_i, \mathbf{x}) + \sigma_n^2$$

Correlation and SHAP Feature Importance

Pearson Correlation

Measures the strength and direction of the linear relationship between an anthropometric biomarker X and age Y :

$$r = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2 \sum_{i=1}^n (y_i - \bar{y})^2}}, \quad -1 \leq r \leq 1$$

SHAP (SHapley Additive exPlanations)

Explains the contribution of each feature to a machine learning prediction:

$$f(\mathbf{x}) = \phi_0 + \sum_{i=1}^M \phi_i$$

where ϕ_0 is the baseline prediction and ϕ_i is the contribution of the i^{th} feature. Feature importance is computed as

$$I_i = \frac{1}{N} \sum_{j=1}^N |\phi_{ij}|,$$

where larger I_i indicates greater importance in predicting age.

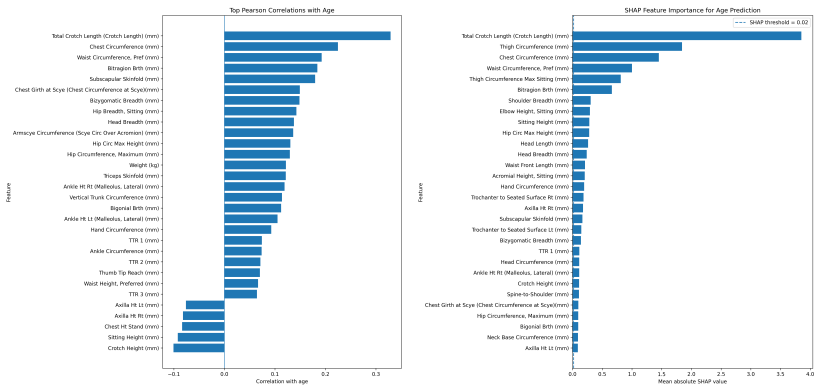


Figure 8: Pearson correlation coefficients (left) and SHAP feature importance scores (right) for anthropometric variables used in age prediction. Correlation quantifies linear feature target relationships, whereas SHAP measures each feature's contribution to the predictive model, accounting for nonlinear effects and feature interactions.

Table 5: Biomarkers showing strong agreement between Pearson correlation and SHAP feature importance for age prediction.

Biomarker	Pearson Rank	SHAP Rank	Interpretation
Total Crotch Length (Crotch Length)	#1	#1	Strongest predictor overall.
Chest Circumference	#2	#3	Strong linear and model importance.
Waist Circumference (Preferred)	#3	#4	Consistently important.
Bitragion Breadth	#4	#6	Important cranial measurement.
Subscapular Skinfold	#5	#18	Linear relationship with moderate nonlinear contribution.
Head Breadth	#10	#12	Consistent anthropometric marker.
Hip Circ Max Height	#11	#10	Good agreement.
Hand Circumference	#19	#16	Moderate importance in both analyses.

Objective

Compare three regression models for Biological age prediction using anthropometric measurements from the CAESAR dataset.

- RBF Support Vector Regression (SVR)
- RBF Kernel Ridge Regression (KRR)
- Gaussian Process Regression (GPR)

Data Preparation

- Age range:18-69 single group
- Stratified train/test split (80% / 20%) within each age group
- Features standardized:

$$x^* = \frac{x - \mu}{\sigma}$$

- Target age standardized during training and converted back after prediction

RMSE Comparison of Weighted-RBF Regression Models

Table 6: Root Mean Square Error (RMSE, years) of the six regression models across a single global group. The models compare correlation-weighted and SHAP-weighted radial basis function (RBF) kernels for Support Vector Regression (SVR), Kernel Ridge Regression (KRR), and Gaussian Process Regression (GPR). The lowest RMSE within each age group is shown in bold.

Age-1 group	Corr-SVR	SHAP-SVR	Corr-KRR	SHAP-KRR	Corr-GPR	SHAP-GPR
Male(1029)	8.25	7.80	8.18	7.84	11.48	11.48
Female (1135)	8.87	8.66	8.87	8.53	12.07	12.07
Combined (2164)	8.51	8.45	8.34	8.12	7.98	7.99

Future Work

- Apply K-means clustering to these 6 models

- Random Forest is built on the base learner of decision trees
- At each node/tree a single marker and a single threshold is chosen. This threshold is chosen to reduce the impurities of the target. This is recurring at each tree until the maximum depth is reached.
- The main parameters for random forest is the number of estimators used, the max depth of tree, and the number of samples used per decision tree

- Use a random search algorithm to test combinations of the many parameters used in RF across a grid of different values selected by the developer
- Train the RF model on 80% of the data to give it enough information to make reliable decisions.
- Pair with K-Means clustering on new data
- Use TreeShap to get marker importance to see how they affect the output of the model
- Run a centile deviation analysis to obtain atypicality ratings for each patient just as in the SVR.

Random Forest Results

Below are the results obtained from testing Random Forest.

Model	Error (yrs)
11 age brackets	11.6
Tuned hyper-parameters	10.7
Bracketed markers	13.3
Bracketed markers + interactions	12.42
Single model, TreeSHAP + K-means	9.09

Table 7: Results from RF tests.

Random Forest Results

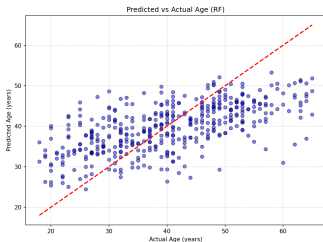


Figure 9: RF with K-Means on single age bracket: RMSE = 9.09 years

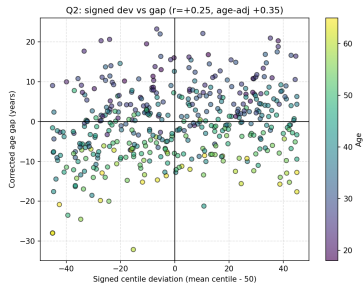


Figure 10: Signed deviation Vs. Age gap

Random Forest Limitations

- Random Forest still underperformed compared to other models even on the larger dataset.
- The piecewise nature of the RF handles other data more elegantly as it is built to withstand harsh interactions and breakpoints. Our data favors more of a linear nature with the continuous trend of aging, indicating the use of Kernel based methods work better with our specific data.

Goals and Plans Moving Forward

Goals moving forward

- Creating a sublevel under Age grouping and understanding the important biomarkers that impact Biological Age.
- Fine tuning the models to be work well with the caesar dataset. (trying different loss functions or different parameters)
- Focusing in on male vs. female predictions.

We would like to thank

- Prof. Peter Wolenski and Dr. Nadejda Drenska for guiding and supporting us.
- Dr. Steven B. Heymsfield from Pennington Biomedical Research Center for giving us the datasets and supporting the evaluation of the models.

References I



Gyaneshwar Agrahari, Kiran Bist, Monika Pandey, Jacob Kapita, Zachary James, Jackson Knox, Steven Heymsfield, Sophia Ramirez, Peter Wolenski, and Nadejda Drenska.

Predicting anthropometric body composition variables using 3d optical imaging and machine learning.

Frontiers in Bioinformatics, 6:1722578, 2026.



Corinna Cortes and Vladimir Vapnik.

Support-vector networks.

Machine Learning, 20:273–297, 1995.



Carlos A. Fermín-Martínez, Alejandro Márquez-Salinas, Enrique C. Guerra, Lilian Zavala-Romero, Neftali Eduardo Antonio-Villa, Luisa Fernández-Chirino, Eduardo Sandoval-Colin, Daphne Abigail Barquera-Guevara, Alejandro Campos Muñoz, Arsenio Vargas-Vázquez, César Daniel Paz-Cabrera, Daniel Ramírez-García, Luis Miguel Gutiérrez-Robledo, and Omar Yaxmehen Bello-Chavolla.

Anthropoage, a novel approach to integrate body composition into the estimation of biological age.

Aging Cell, 22(1):e13756, 2023.

References II



J. A. K. Suykens and J. Vandewalle.

Least squares support vector machine classifiers.

Neural Processing Letters, 9:293–300, 1999.

Thank You!!
Questions?