

Wearable Sweat Sensor Data Analysis: Artifact Detection and Imputation

Luca Guimont, lg5g20@soton.ac.uk
University of Southampton, Department of Vision Learning & Control

Introduction

Water and sodium are lost from the body during exercise, primarily via sweat. Improperly replenishing these losses can lead to dehydration, overhydration, or sodium imbalances, all of which negatively impact performance. Optimizing an athlete's hydration is therefore a key factor in improving their performance. To this end, FLOWBIO produces a wearable sweat sensor (the S1) to measure sweat sodium concentration and fluid loss during exercise, which is used in conjunction with a third-party smart watch and the FLOWBIO app to produce recommendations of how much fluid and sodium an athlete should consume before and after exercise. Wearable sensors are often susceptible to artifacting and noise from external sources. In the case of the S1, moving the device from its affixed reading position on the arm or back can break the continuous collection of sweat, causing the output to become inaccurate until sweat-saturation is restored. The challenge is to accurately detect the start and end of these events, then to impute data over the gap that mimics the real underlying sweat patterns. This research contributes a means to detect and diagnose data artifacts, as well as an analysis of imputation methods applied to the FLOWBIO exercise sessions dataset. These preliminary contributions lay the groundwork for more sophisticated imputation and prediction techniques to improve FLOWBIO's services.

Methods

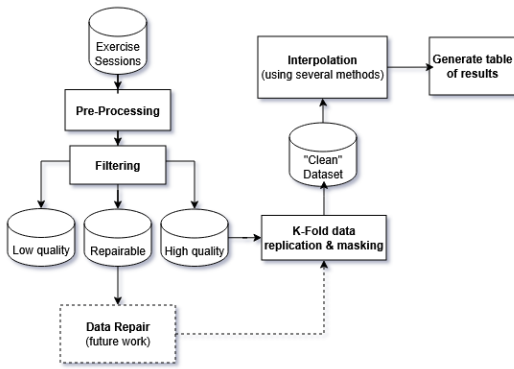


Figure 1: Proposed interpolation experiment

Additional features were engineered from the base sensor readings, such as differences per timestep. Each timestep of a session was classified by data quality using a series of filters over the features. The thresholds of these filters were tuned by comparing the filter-generated classifications to manually classified examples. Data was considered 'high-quality' when the sensor was working normally. Some sections were classified as 'repairable' where the data pattern matched a known artifacting source that non-linearly scales the true data signal, rather than overwrit-

ing it. Data was classified 'low-quality' only during events where the sensor was unable to properly measure the athlete's sweat. These low-quality sections are the primary targets for interpolation. However, determining the best interpolation method to use over the S1 data could not be done using the real low-quality sections as there is no ground-truth to compare to. Instead, a 'clean' dataset was made by compiling and normalizing sufficiently long, uninterrupted sections of high-quality sensor data. Over 12,000 sections were extracted this way. These high-quality sections were then replicated in a K-fold manner, each copy having a unique sub-section masked to emulate some low-quality data being removed. A battery of interpolation methods was then used to fill in the artificially missing data, taking the difference between the interpolated data and the original un-masked data at each timestep as the error and calculating RMSE.

Results

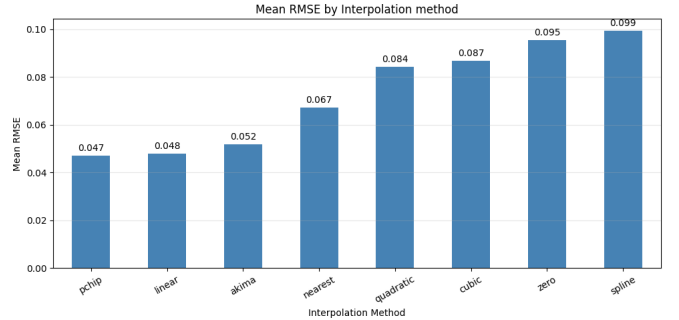


Figure 2: Interpolation error by method

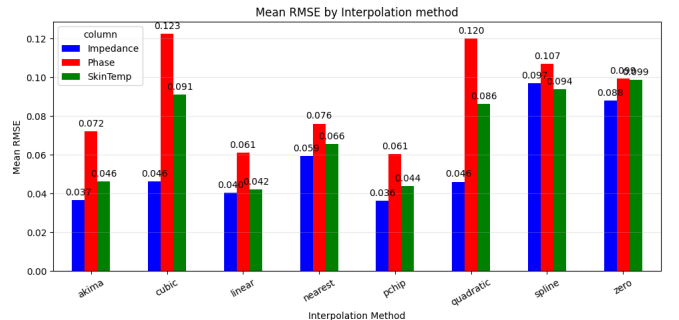


Figure 3: Interpolation error by method, split by sensor reading

The results of the experiment show that pchip marginally outperformed linear with, Akima closely following. Nearest value, quadratic, and cubic fared worse but still outperformed forward-fill (labeled as 'zero' as a zero-order spline was used). The order-3 spline performed the worst among the methods tested. Interestingly, the sensor reading Phase produced the most error, suggesting a more non-linear structure than impedance or skin-temp.