

SISTEMAS DE RADAR INDOOR PARA A IDENTIFICAÇÃO DA IDENTIDADE DE DOIS SUJEITOS ATRAVÉS DE INTELIGÊNCIA ARTIFICIAL

DAVID PEREIRA¹; ANDRÉ ROUCO²; DANIEL ALBUQUERQUE³

¹Mestrando em Engenharia Biomédica, Universidade de Aveiro, Portugal, david.a.mpereira@ua.pt.

²Instituto de Telecomunicações, DETI, Universidade de Aveiro, Portugal, andrerouco@ua.pt.

³Instituto de Telecomunicações, ESTGA, Universidade de Aveiro, Portugal, dfa@ua.pt.

RESUMO

Este projeto explora a utilização de sistemas de radar indoor para distinguir entre dois indivíduos através do estudo da sua “assinatura” Radar, mais especificamente o *Heart Rate* (HR), o *Breath Rate* (BR), os valores de posição (X, Y e Z) do indivíduo e os valores de doppler e SNR. Os dados foram recolhidos durante 10 dias e a 4 momentos do dia (manhã, almoço, tarde e noite), com 10 minutos por sessão. Foram feitos 2 testes, sendo que no primeiro foram escolhidos os dados teste e de treino aleatoriamente, onde foi obtida uma precisão média de 99.4% ao longo dos 4 algoritmos de *machine learning* utilizados, e no segundo os algoritmos foram testados com os dados adquiridos no período da manhã e treinados com os restantes, de seguida, foram testados com os dados do período do almoço e treinados com os restantes e assim sucessivamente, onde foi obtida uma precisão média de 85.9%.

Palavras-chave: Sistemas de radar indoor; *Heart Rate* (HR); *Breath Rate* (BR); *Machine Learning* (ML), identificação biométrica (ID).

INDOOR RADAR SYSTEMS FOR IDENTIFYING THE IDENTITY OF TWO SUBJECTS USING ARTIFICIAL INTELLIGENCE

ABSTRACT

This project explores the use of indoor radar systems to distinguish between two individuals by studying their Radar ‘signature’, more specifically the Heart Rate (HR), the Breath Rate (BR), the position values (X, Y and Z) of the individual and the doppler and SNR values. The data was collected over 10 days and at 4 times of the day (morning, lunch, afternoon and evening), with 10 minutes per session. Two tests were carried out: in the first, the test and training data were chosen at random and an average accuracy of 99.4% was obtained across the four machine learning algorithms used; in the second, the algorithms were tested with the data acquired in the morning and trained with the rest, then tested with the lunchtime data and trained with the rest, and so on, where an average accuracy of 85.9% was obtained.

Keywords: Indoor radar systems; Heart Rate (HR); Breath Rate (BR); Machine Learning (ML); biometric identification (ID).

INTRODUÇÃO

A confirmação da identidade de uma pessoa passa normalmente pela verificação de uma entidade concreta relacionada com essa pessoa. Normalmente, uma entidade pode envolver: a posse de uma pessoa, por exemplo, uma chave; o conhecimento de uma pessoa, como uma palavra-passe; ou mesmo uma combinação de ambos, por exemplo, caixas multibanco que utilizam um cartão e um pin para estabelecer uma identidade. No entanto, a utilização de bens como meio de identidade tem os seus inconvenientes, uma vez que o utilizador pode perder os tokens ou alguém não autorizado pode obtê-los e abusar dos

privilégios do utilizador autorizado. A utilização de conhecimentos, por outro lado, pode ser difícil de recordar ou ser facilmente adivinhada por outros. Por conseguinte, é imperativo utilizar um meio de verificação mais seguro: a biometria surgiu como uma solução possível (JAIN; BOLLE; PANKANTI, 1999).

A biometria define o reconhecimento de uma pessoa com base nas suas características fisiológicas ou comportamentais (LIM; GUPTA, 2003). Estas não podem ser perdidas ou esquecidas e representam uma componente tangível de algo que se é, o que as torna vantajosas como método de reconhecimento (JAIN; BOLLE; PANKANTI, 1999).

Métodos como a impressão digital ou o reconhecimento facial têm uma utilização generalizada, no entanto, há uma necessidade crescente de explorar outras modalidades que sejam menos vulneráveis a ataques furtivos ou de falsificação (KUMAR et al., 2019).

Deste modo, neste projeto foram usados sistemas de radar *indoor* para que, com a medição do batimento cardíaco, do ritmo respiratório, da variação da posição, do valor de *Doppler* e *Signal Noise Ratio (SNR)* de dois indivíduos distintos, fosse possível distingui-los com recurso a algoritmos de machine learning.

Os principais objetivos deste trabalho são:

- Verificar se, com apenas os sinais de Radar, é possível, com uma boa precisão, distinguir um indivíduo A de um indivíduo B;
- Avaliar aplicações do dia a dia de um algoritmo deste género;
- Avaliar a estabilidade dos sinais biométricos ao longo do tempo, recorrendo a aquisição de sinais em diferentes dias e em diferentes alturas do dia.

METODOLOGIA

O método utilizado para a recolha dos dados pode ser observado na Figura 1. O sujeito foi sentado numa cadeira a uma distância de 1 metro do radar. Este radar encontra-se a um metro do chão e com uma inclinação de 0° de frente para o utilizador. Esta montagem foi baseada nas instruções apresentadas em (TEXAS INSTRUMENTS, 2024).

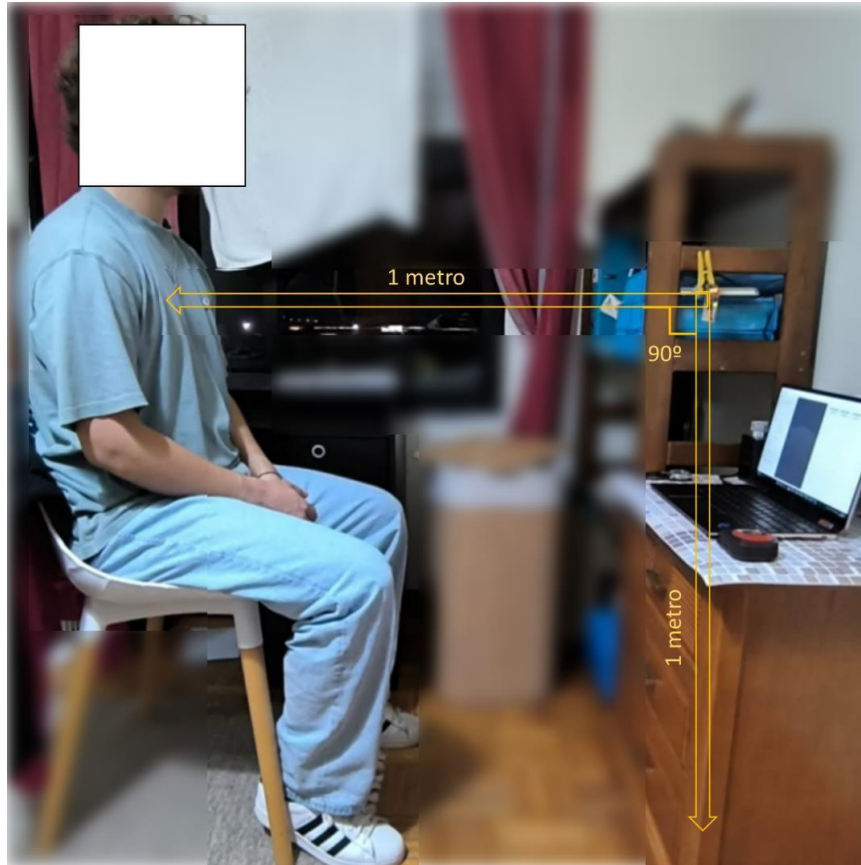


Figura 1 – Montagem realizada para a recolha de dados. Fonte: acervo dos autores, 2024.

Foram feitas sessões de recolha de dados a dois participantes em 4 alturas diferentes do dia, sendo uma de manhã, uma imediatamente após o almoço, uma ao fim da tarde e uma de noite de modo a aumentar a variabilidade dos dados.

Ambos os participantes assinaram um termo de consentimento informado para utilização dos seus sinais vitais para este estudo. Para além disso, estes testes foram também aprovados pela comissão ética da Universidade de Aveiro.

Os sinais foram gravados com a duração de 15 minutos, onde foram eliminados os primeiros 3 e os últimos 2 minutos para evitar, no primeiro caso, a existência de erros de ajuste provenientes do radar, visto que este demora algum tempo a detetar e estabilizar os sinais do sujeito, e no último caso, garantir que todas as amostras retiradas tinham o mesmo tempo. Posteriormente os sinais que ficaram com a duração de 10 minutos foram ainda divididos em 10 sinais de 1 minuto e separados nas variáveis X, Y e Z, *doppler*, *SNR*, *noise* e *Track Index* presentes na lista *pointCloud* fornecida pelo sistema de radar e as variáveis *id*, *rangeBin*, *breathDeviation*, *heartRate*, *breathRate*, *heartWaveform* e *breathWaveform* presentes na lista *vitalsDict* fornecida também ela pelo sistema de radar.

Recorrendo ao código fornecido pela TI e por motivos de simplicidade, o código fornecido foi modificado de modo que, enquanto o programa apresenta os sinais de

HeartWaveform e *BreathWaveform* em tempo real, este guarde os valores numa nova lista e, no final, os guarde num ficheiro de texto. Posteriormente, estes dados são carregados para um ficheiro de código para que se possa proceder ao tratamento dos mesmos. Como já referido anteriormente, os dados consistem nas matrizes de *vitalsDict* e *pointCloud* que contêm as variáveis referentes aos sinais vitais e à posição, respetivamente.

Inicialmente, com o objetivo de realizar o pré-processamento dos sinais vitais, os dados foram carregados para um ficheiro de *Python*. De seguida foi-lhes aplicado uma função que, dados o sinal obtido experimentalmente e o tempo de aquisição de dados (15 minutos), retornava 5 listas onde se encontram os valores de *BreathDeviation*, *HeartRate*, *BreathRate*, *HeartWaveform* e *BreathWaveform*.

Esta função realiza as tarefas de, para além de separar as variáveis, eliminar os primeiros 3 e últimos 2 minutos por motivos abordados anteriormente e separar os restantes dados em segmentos de 1 minuto.

Relativamente aos sinais de posição, foi-lhes aplicado uma função semelhante que, ao receber o sinal e o tempo de aquisição nos retorna 6 listas, sendo elas a *xPoints*, a *yPoints*, a *zPoints*, a *dopplerPoints*, a *snrPoints* e a *noisePoints*. Além disso, esta função também retorna a variável *timeBetweenFrames* que vai ser necessária para o cálculo das velocidades.

As tarefas realizadas por esta última função são idênticas às realizadas pela primeira, a única diferença entre ambas é que para os sinais vitais foi necessário realizar um passo para correção de valores escritos em notação científica para que estes sejam reconhecidos pelo programa de *python* como números escritos em notação científica.

Na Figura 2 é possível observar a representação gráfica da *HeartWaveform* e da *BreathWaveform* para o sujeito 1 e para o sujeito 2.

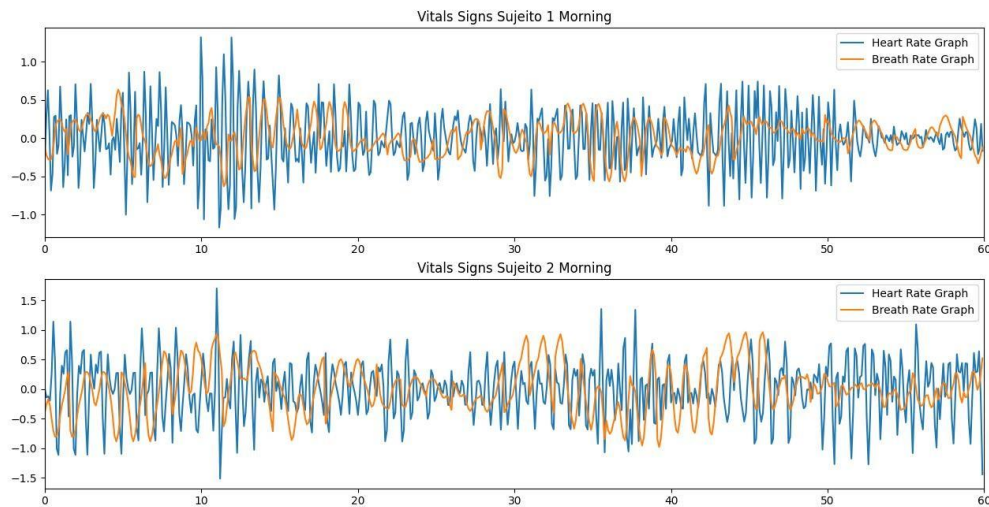


Figura 2 – Representação gráfica da *HeartWaveform* e da *BreathWaveform* do sujeito 1 e do sujeito 2. Fonte: acervo dos autores, 2024.

De seguida, devido à natureza do radar e recorrendo às matrizes *xPoints*, *yPoints* e *zPoints*, foi calculada a velocidade utilizando dois métodos distintos para cada um dos 3 eixos. No primeiro método, a velocidade é calculada através da equação (1), onde $(x_{n-4} - x_n)$ corresponde à distância entre o primeiro ponto e o ponto 5 índices à sua frente e Δt corresponde à diferença temporal entre esses dois pontos.

$$v_1 = \frac{(x_{n-4} - x_n)}{\Delta t}, n = 5, 10, 15, \dots \quad (1)$$

Por outro lado, no segundo método a velocidade é calculada pela equação (2) onde μ_1 corresponde à média dos 5 primeiros pontos, μ_2 corresponde à média dos 5 pontos subsequentes e Δt corresponde à diferença temporal entre o ponto médio de cada conjunto de pontos.

$$v_2 = \frac{(\mu_2 - \mu_1)}{\Delta t} \quad (2)$$

O cálculo da velocidade utilizando dois métodos diferentes é útil para lidar com a presença de outliers nos dados de posição. A fórmula 1 é muito mais afetada por *outliers*, embora, na ausência destes, o valor de velocidade obtido seja mais preciso. Em contraste, a fórmula 2 é menos suscetível a outliers devido à utilização de médias ao invés de pontos, mas produz resultados de velocidade menos realistas. Para o cálculo de *features* foram utilizados ambos os valores de velocidade.

Já com o pré-processamento dos dados finalizado, passamos agora à extração de *features*. O uso de *features* tem como principal objetivo caracterizar os dados numa representação que ao mesmo tempo reduz os efeitos de ruídos e variabilidade intra sujeito e realça as diferenças entre sujeitos (KUMAR et al., 2018).

Com base em estudos semelhantes já realizados (BASHEER; WANI; KUMAR, 2020) (PAIVA et al., 2021), é possível afirmar que os melhores resultados obtidos, usualmente, utilizam a morfologia do sinal HR e BR como principais intervenientes na distinção entre sujeitos. Tendo esta informação como base foi calculado, para cada conjunto de dados disponíveis, a média do seu valor absoluto, a sua potência, a variância, o desvio padrão, a *kurtosis* e a *skewness*.

Para além disso, exclusivamente para os sinais de *HeartWaveform* e *BreathWaveform*, foi calculado também a distância média entre picos e, após visualização do *boxplot* (Figura 3) e posterior eliminação dos outliers, a amplitude máxima dos sinais.

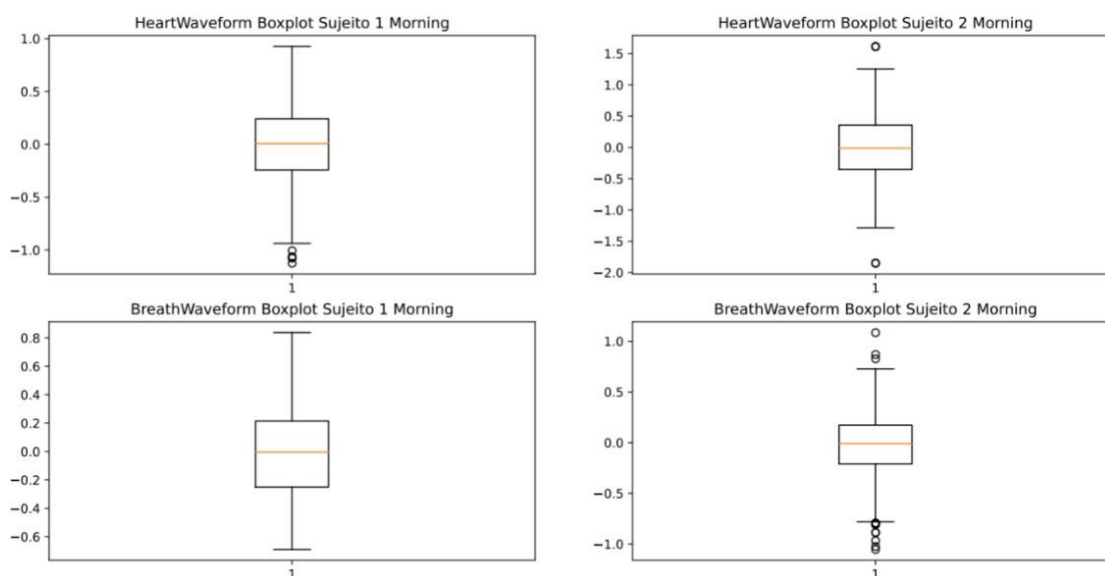


Figura 3 – Boxplot da HeartWaveform e da BreathWaveform do sujeito 1 e do sujeito 2 para visualização dos outliers. Fonte: acervo dos autores, 2024.

Para uma representatividade ainda maior dos sinais cardíaco e respiratório, realizou-se a *Fast Fourier Transform* (FFT), apresentada na Figura 4, sendo possível calcular todas as variáveis previamente calculadas para os restantes sinais à exceção da distância média entre picos, a amplitude máxima dos sinais, a *kurtosis* e a *skewness*, visto que neste contexto não fariam sentido.

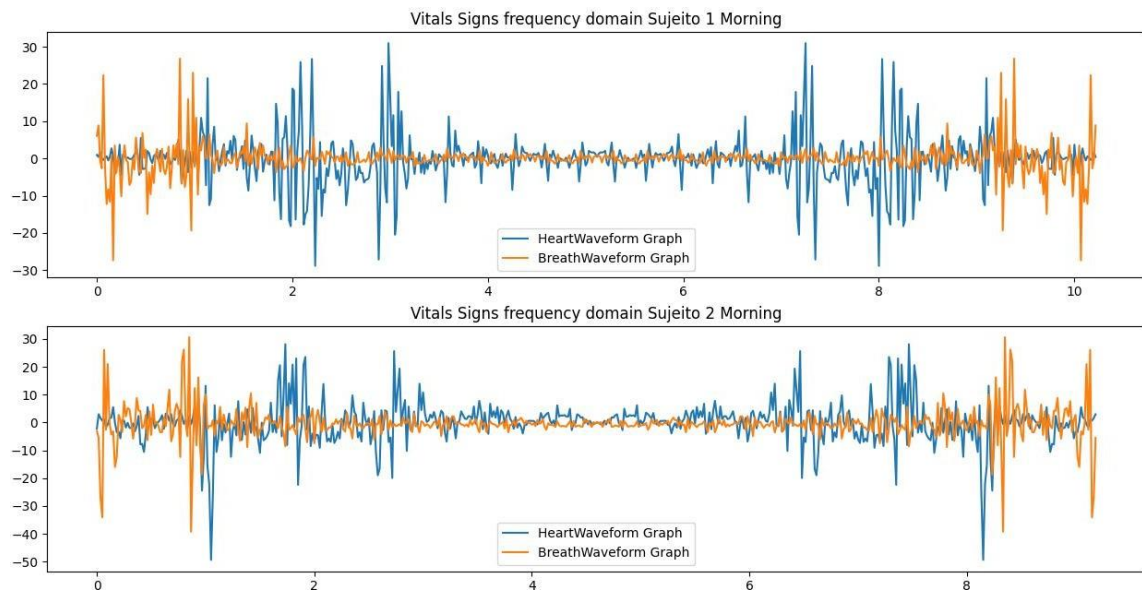


Figura 4 – Representação gráfica da FFT da HeartWaveform e da BreathWaveform do sujeito 1 e do sujeito 2. Fonte: acervo dos autores, 2024.

Ainda foi calculada também para as FFTs a potência do sinal para cada espaço de 0.1 Hz, ou seja, o sinal foi dividido em segmentos de 0.1 Hz e, para cada um desses segmentos, foi calculada a energia do sinal.

Esta análise resultou num total de 170 *features* para cada segmento de um minuto. Estas *features* podem ser encontradas na Tabela 1 apresentada abaixo, onde existem 11 grupos de dados que apenas são calculadas as *features* apresentadas na tabela para a variável Data e 2 grupos de dados, que são os gráficos de HR e BR aos quais são calculadas as *features* apresentadas na tabela para a variável WaveformData.

Tabela 1 – Tabela de features

Data	WaveformData
Média do valor absoluto	Média do valor absoluto
Energia do sinal	Energia do sinal
Variância	Variância
Devio Padrão	Devio Padrão
Kurtosis	Kurtosis
Skewness	Skewness
	Distância entre picos
	Amplitude
	Média do valor absoluto da FFT
	Potência da FFT por 0.1 Hz
	Potência total da FFT
	Variância da FFT
	Desvio padrão da FFT

Fonte: acervo dos autores, 2024.

RESULTADOS E DISCUSSÃO

Num primeiro teste os dados foram escolhidos de forma aleatória através do uso da função *train_test_split* fornecida pela biblioteca de *python sklearn*. Esta função escolhe uma percentagem de dados aleatórios para os dados de teste e de treino de acordo com a percentagem definida pelo utilizador. Neste caso, devido à quantidade de dados disponível, foi usada uma percentagem para os dados de teste de 30% e 70% para os dados de treino.

Antes de passar à fase dos resultados, é importante realçar a necessidade de normalizar os dados antes de os usar para treinar e testar os algoritmos, visto que, sem a normalização dos mesmos, os algoritmos de machine learning podem ter um mau desempenho se as *features* não se assemelharem a dados normalmente distribuídos (por exemplo, *Gaussianos* com média 0 e variância unitária).

Assim foi aplicado o *StandardScaler* disponível na biblioteca de *python sklearn* que, basicamente, aplica a fórmula 3 ao vetor das *features*, onde μ representa a média das *features* de treino e σ representa o desvio padrão. Esta fórmula transforma o vetor das features para que se assemelhe a uma distribuição normal, com média 0 e variância unitária.

Para melhor observar os resultados deste teste foi feito um *loop* onde o código corre 15 vezes e, entre cada iteração, são escolhidos novos dados de treino e teste. No final, o resultado de precisão de cada algoritmo é calculado através da média da precisão das 15 iterações acompanhado do desvio padrão das mesmas.

Através deste teste foi possível obter resultados bastante satisfatórios. Para o algoritmo KNN foi obtida uma precisão de $97.5 \pm 3.1\%$ e para os algoritmos SVM, LDA e DT a precisão foi de $100.0 \pm 0.0\%$.

Para ter uma melhor percepção dos valores de precisão em cada uma das iterações foi elaborado um gráfico que mostra os valores de precisão dos diferentes algoritmos para cada conjunto de dados de teste e treino, como é possível observar na Figura 5.

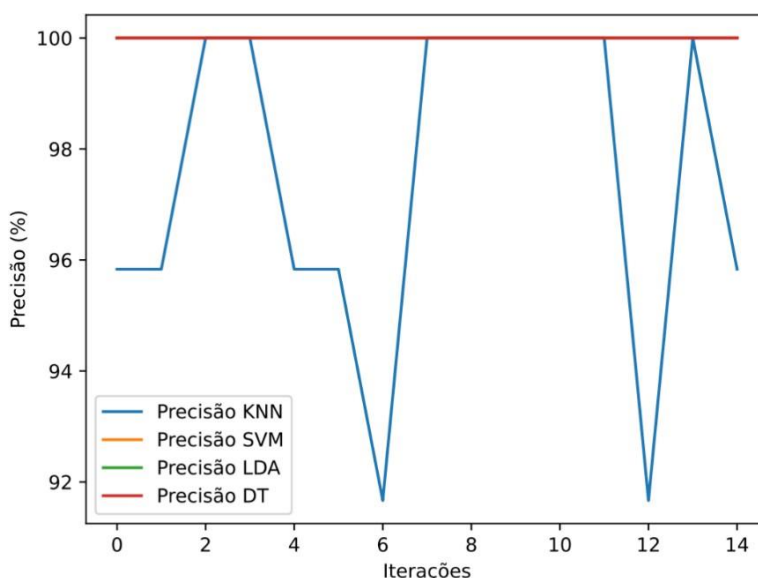


Figura 5 – Valores de precisão para cada algoritmo ao longo das 15 iterações. Fonte: acervo dos autores, 2024.

Tendo em conta os resultados obtidos, podemos afirmar que os valores de precisão foram bastante satisfatórios visto que, para todos os algoritmos utilizados, o valor da precisão médio foi acima dos 97%.

Para além disso, é também possível comparar o desempenho dos vários algoritmos utilizados, sendo notável que os algoritmos SVM, LDA e DT são os mais precisos na classificação do sujeito, não tendo errado nenhuma das suas classificações ao longo das 15 iterações.

Por outro lado, o algoritmo KNN mostrou desempenho também excelente embora não perfeito. Isto sugere que todos os algoritmos são adequados para este tipo de classificação,

embora a escolha final deva considerar outros fatores como a complexidade computacional e o tempo de execução. É também de realçar que os desvios padrões são nulos ou extremamente baixos, o que mostra excelente estabilidade dos algoritmos variabilidade dos dados.

No segundo teste, ao contrário do que foi realizado para o primeiro, foram escolhidos manualmente os dados para treino e teste dos 4 diferentes algoritmos. Desta forma e com o objetivo de testar a variabilidade de momento para momento do dia, os algoritmos foram testados com dados do mesmo momento do dia e treinados com os restantes, ou seja, numa primeira iteração, os dados de teste vão corresponder aos dados adquiridos durante a manhã, enquanto que os de treino serão os restantes, numa segunda iteração, os dados de teste vão corresponder aos dados adquiridos durante a hora de almoço, enquanto que os de treino serão os restantes e assim sucessivamente.

Desta forma é possível testar a variabilidade de período para período do dia, sendo expectável que esta exista, mas tenha baixa influência nos valores da precisão (REFINETTI, 2016). De referir que foi utilizado o mesmo método de normalização utilizado para o teste 1.

Utilizando os dados da manhã como dados de teste e os restantes dados disponíveis como dados de treino foi possível obter uma precisão de 90% para o algoritmo KNN e SVM, apresentando 2 casos onde erraram a sua classificação como é possível observar na matriz de confusão presente na Figura 6. Para o algoritmo LDA a sua precisão foi de 85% apresentando um total de 3 erros de classificação. Por último temos o algoritmo DT, que apresentou uma precisão de 100%, tendo acertado todas as suas classificações.

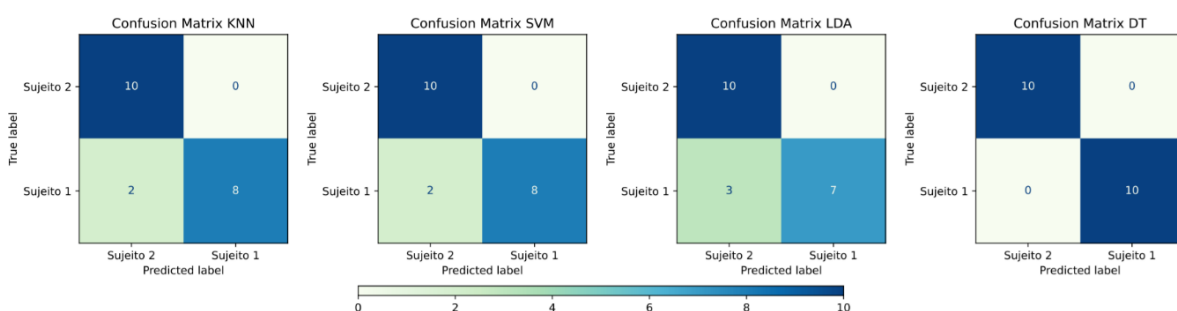


Figura 6 –Matrizes de Confusão de cada algoritmo para os dados de teste adquiridos de manhã..
Fonte: acervo dos autores, 2024.

Para o caso em que os dados de teste utilizados foram os adquiridos à hora de almoço e os restantes foram utilizados como dados de treino, os valores de precisão foram relativamente parecidos. Assim, a precisão calculada para o algoritmo KNN foi de 95%, com um erro de classificação. Por outro lado, o algoritmo de SVM apresentou uma precisão mais baixa de 70%, tendo um total de 6 erros de classificação. Por último, tanto o algoritmo de DT

como o LDA apresentaram uma precisão de 100% tendo acertado todas a classificações realizadas, tal como é possível observar na Figura 7.

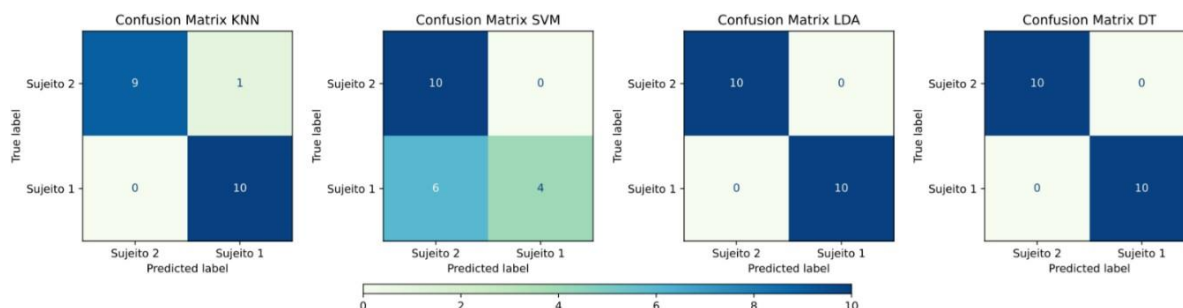


Figura 7 – Matrizes de Confusão de cada algoritmo para os dados de teste adquiridos no período de almoço. Fonte: acervo dos autores, 2024.

Para os dados adquiridos ao final da tarde como dados de teste e os restantes disponíveis como dados de treino foram obtidos números de precisão maioritariamente menos satisfatórios. Desta forma, a precisão obtida para os algoritmos de KNN foi de 55%, tendo errado a sua classificação 9 vezes. Para o caso do algoritmo SVM, este apresentou uma precisão de 100%, tendo acertado todas as classificações. O algoritmo LDA apresentou uma precisão de 65%, o que resulta num total de classificações erradas de 7. Por último o algoritmo DT obteve uma precisão de 50%, classificando todos os dados como pertencentes ao sujeito 1, tal como é possível observar na Figura 8. É de notar que o algoritmo DT classificou todos como Sujeito 1, o que pode significar que os dados adquiridos no período do final da tarde do sujeito 2 são relativamente distintos dos restantes, como é possível observar no gráfico da Figura 9.

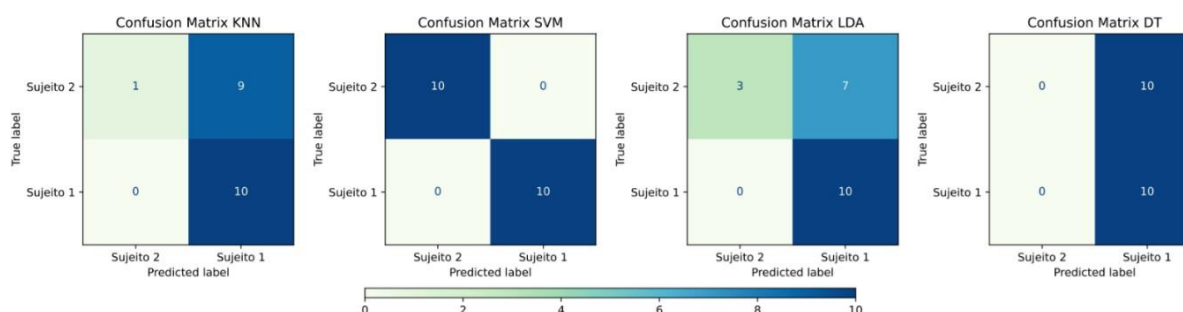


Figura 8 – Matrizes de Confusão de cada algoritmo para os dados de teste adquiridos ao final da tarde. Fonte: acervo dos autores, 2024.

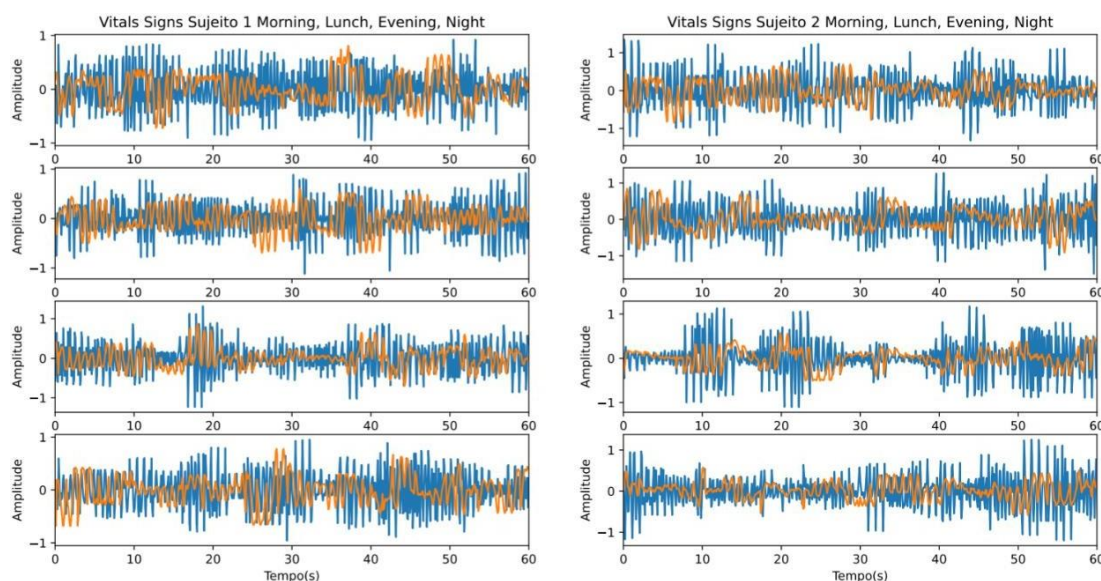


Figura 9 – Sinais vitais do sujeito 1 e sujeito 2 ao longo dos diferentes momentos do dia (Manhã, Almoço, final de tarde e noite respectivamente). Fonte: acervo dos autores, 2024.

Por último, ao utilizar os dados adquiridos à noite como dados de teste e os restantes disponíveis como dados de treino foi possível obter números de precisão bastante mais satisfatórios. Desta forma, a precisão obtida para os algoritmos de KNN, DT e LDA foi de 100%, tendo acertado todas as classificações realizadas. Por outro lado, o do algoritmo SVM apresentou uma precisão de 75%, tendo errado um total de 5 classificações, tal como é possível observar na Figura 10.

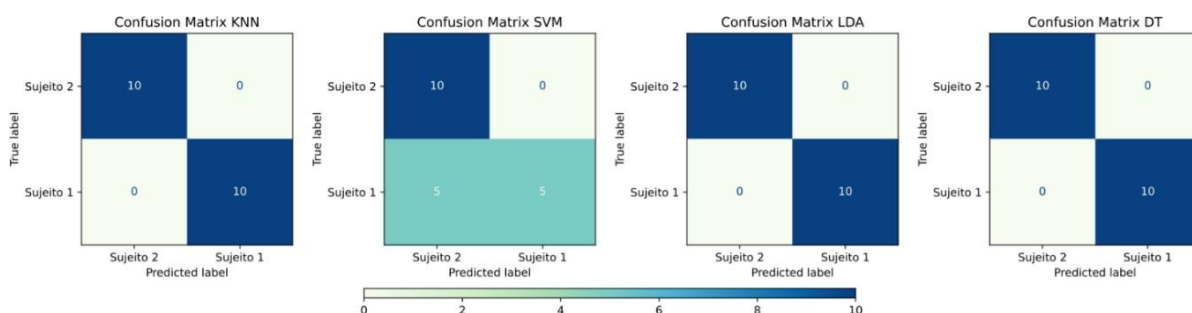


Figura 10 – Matrizes de Confusão de cada algoritmo para os dados de teste adquiridos de noite. Fonte: acervo dos autores, 2024.

Com estes dados também é possível calcular a precisão média para cada algoritmo. Assim, a precisão para o algoritmo KNN foi $85.0 \pm 20.4\%$, a precisão do algoritmo SVM foi de $83.8 \pm 13.8\%$, para o algoritmo LDA a precisão foi $87.5 \pm 16.6\%$ e por último, o valor de precisão para o algoritmo DT foi $87.5 \pm 25.0\%$.

Após observação dos resultados obtidos com este teste e visto que os valores de precisão são baixos e relativamente distintos para cada conjunto de dados de teste, é possível afirmar que existe variabilidade dos sinais biométricos ao longo dos diferentes momentos do dia embora esta não seja suficientemente significativa para baixar consideravelmente os resultados de precisão do teste 1 referido anteriormente. Esta conclusão já era expectável pois estudos anteriores, (REFINETTI, 2016) (NICKLE et al., 2015), indicam que os sinais biométricos, embora sujeitos a variações naturais ao longo do dia devido a fatores como níveis de atividade, alimentação e estados emocionais, tendem a manter características consistentes que permitem uma classificação precisa.

Tendo em atenção agora os valores do desvio padrão, podemos observar que o algoritmo SVM é o único que apresenta um desvio menor do que 15%, o que mostra que este algoritmo é o mais estável para valores de treino com menor variabilidade. Na outra ponta do espectro temos o algoritmo DT que, para o teste 1 apresenta uma precisão de 100% e no teste 2 apenas 87.5%, onde para um dos dados de teste escolhidos classifica erradamente todos os dados como pertencentes ao mesmo sujeito. Globalmente, os valores de desvio padrão deste teste são extremamente maiores do que os do teste 1, o que significa que quanto maior a variabilidade presente nos teste de treino, melhores resultados vão ser apresentados pelos algoritmos.

CONSIDERAÇÕES FINAIS

Durante este projeto, foi testada a hipótese da utilização de vários algoritmos de *machine learning* na classificação da identidade de sujeitos utilizando sinais vitais adquiridos com recurso a um sistema de radar.

Este projeto tem um contributo considerável para a área de biometria ao demonstrar a viabilidade do uso de sistemas de radar combinados com algoritmos de *machine learning* para a identificação de sujeitos. Através de testes empíricos, foi possível não apenas calcular a precisão dos métodos, mas também fornecer conhecimento acerca da sua aplicabilidade prática.

No primeiro teste foi possível avaliar este tipo de algoritmos num contexto mais real, visto que, para o algoritmo ser válido e útil para futuras aplicações, este tem de ser capaz de identificar o sujeito pelos seus sinais vitais independentemente do momento e do dia.

Já para o segundo teste explorado, o principal objetivo foi testar a variabilidade dos sinais vitais ao longo de vários dias e vários momentos do dia. Observando os resultados obtidos foi possível concluir que, como já discutido previamente, estes variam, embora não o suficiente para afetar consideravelmente os resultados do primeiro teste.

Tendo em conta os resultados de ambos os métodos é possível afirmar que, embora exista variabilidade nos dados, esta não é suficientemente elevada para afetar significativamente o resultado das precisões do segundo método, provando assim que os sinais adquiridos com o sistema de RADAR são sinais fiáveis para fins de biometria, podendo assim ser considerados válidos para futuras tarefas de identificação e autenticação.

Embora os resultados de ambos os métodos tenham sido bastante positivos, principalmente os do teste 1, é essencial considerar que este tem algumas limitações. A principal limitação deste projeto seria a sua base de dados, visto que esta é relativamente diminuta (apenas 40 minutos de dados para cada sujeito) e pouco variável (os dados foram gravados no espaço de aproximadamente 10 dias e apenas 4 momentos do dia), o que pode dificultar a generalização dos resultados a outros contextos e populações. Futuras investigações na área devem considerar a possibilidade de ter maior volume e variabilidade de dados, além de mais participantes e de várias origens. Também em futuras investigações deve ser explorada a utilização de diferentes algoritmos de *machine learning* tal como melhorar os que foram neste projeto utilizados, através da otimização dos *hyperparameters*. Além disso, integrar métodos de fusão de dados que combinam ainda mais sinais biométricos, como por exemplo a temperatura corporal com recurso a câmaras térmicas, pode aumentar ainda mais a precisão e a fiabilidade dos sistemas de identificação.

Por último, também seria proveitoso a aplicação destes algoritmos em tempo real e em ambientes diversificados. Testes em condições variáveis, como interferências de sinal e movimentos, são essenciais para garantir que os sistemas sejam adaptáveis e eficientes em situações do mundo real.

Em resumo, os resultados deste estudo são encorajadores e indicam que a integração de sistemas de radar com algoritmos de *machine learning* podem ser uma estratégia eficaz para o uso de sinais vitais como meio de identificação e até autenticação. A seleção do algoritmo mais adequado deve levar em consideração não apenas a precisão, mas também a complexidade e os custos computacionais.

Os avanços tecnológicos em sensores e algoritmos de *machine learning* abrem caminhos promissores para melhorar a segurança e a eficiência em diversas aplicações biométricas. A robustez de algoritmos como o SVM evidencia que é viável alcançar elevados

níveis de precisão, o que é essencial para aplicações em áreas como segurança, saúde e monitorização.

No entanto também é necessário considerar as consequências éticas e sociais associadas a este tipo de dispositivos. Devido à facilidade de obtenção e à singularidade deste tipo de dados, estes requerem rigorosos protocolos de privacidade e segurança de maneira a proteger os dados dos utilizadores. A transparência nas práticas de aquisição e uso de dados e a conformidade com regulamentações de proteção de dados são essenciais para garantir o uso responsável da tecnologia.

REFERÊNCIAS

JAIN, Anil K.; BOLLE, Ruud; PANKANTI, Sharath. **Biometrics: Personal Identification in Networked Society**. Boston: Kluwer Academic Publishers, 1999. Disponível em: <https://pdfroom.com/books/biometrics-personal-identification-in-networked-society/9qlgyRGq2MG>. Acesso em: 9 out. 2024.

LIM, J. S.; GUPTA, M. **Biometric Authentication: A Review of the State-of-the-Art**. In: 2003 IEEE International Conference on Information Technology: Research and Education (ITRE), 2003. p. 272-277. Disponível em: <https://ieeexplore.ieee.org/document/1193209>. Acesso em: 9 out. 2024.

KUMAR, R. et al. **Evolution, current challenges, and future possibilities in ECG biometrics**. *Journal of Medical Systems*, v. 43, n. 10, p. 1-14, 2019. Disponível em: <https://link.springer.com/article/10.1007/s10916-019-1492-7>. Acesso em: 9 out. 2024.

TEXAS INSTRUMENTS. **Radar Toolbox**. Disponível em: https://dev.ti.com/tirex/explore/node?node=A__AMK61TkbsnY89Axyqsuxmw_radar_toolbox_1AslXXD__LATEST. Acesso em: 9 out. 2024.

KUMAR, R. et al. **Evolution, Current Challenges, and Future Possibilities in ECG Biometrics**. 2018 *IEEE International Conference on Bioinformatics and Biomedicine (BIBM)*. p. 157-164. Disponível em: https://www.researchgate.net/publication/325934871_Evolution_Current_Challenges_and_Future_Possibilities_in_ECG_Biometrics. Acesso em: 9 out. 2024.

BASHEER, M. F.; WANI, A. U.; KUMAR, R. **ELM and K-NN Machine Learning in Classification of Breath Sounds Signals**. 2020 *IEEE International Conference on Bioinformatics and Biomedicine (BIBM)*. p. 1-6. Disponível em: https://www.researchgate.net/publication/340088530_ELM_and_K-nn_machine_learning_in_classification_of_breath_sounds_signals. Acesso em: 9 out. 2024.

PAIVA, L. R.; ARAGÃO, C. F.; PACHECO, F. M.; DA SILVA, A. R.; DE OLIVEIRA, J. C. **A Hybrid Model for Gait Recognition Based on Pressure Sensors Data**. *Sensors*, v. 21, n. 9, p. 3172, 2021. Disponível em: <https://www.mdpi.com/1424-8220/21/9/3172>. Acesso em: 9 out. 2024.

REFINETTI, R. **Circadian Physiology**. Boca Raton: CRC Press, 2016.

NICKLE, D. et al. **A Comparative Study of Machine Learning Algorithms for Biometrics: A Case Study in Gender Recognition**. In: *2015 IEEE International Conference on Biometrics: Theory, Applications and Systems (BTAS)*. p. 1-6. Disponível em: <https://ieeexplore.ieee.org/document/7273870>. Acesso em: 9 out. 2024.