

國立臺灣大學電機資訊學院生醫電子與資訊學研究所

碩士論文

Graduate Institute of Biomedical Electronics and Bioinformatics

College of Electrical Engineering and Computer Science

National Taiwan University

Master's Thesis

腦電圖於預測疑似腦炎患者自體抗體之可行性：

先導機器學習研究

Evaluating Electroencephalography for Predicting Autoantibody
in Suspected Encephalitis: A Pilot Machine-Learning Study

鄭章佑

Chang-Yu Cheng

指導教授：陳志宏 博士

Advisor: Jyh-Horng Chen, Ph.D.

中華民國 114 年 6 月

June 2025



誌謝



能夠完成這個前導研究，要特別感謝我的指導教授陳志宏老師、臺大醫院饒敦醫師、實驗室的學長姐以及學弟妹對我的支持與鼓勵，尤其是虹誼學姊、子毅學長、宜勁學長、唯絃學弟。也要感謝在研究所各位講課的老師的指導，以及碩士論文口試委員(臺大生理所郭鐘金教授、臺大生醫電資所魏安琪副教授、臺大醫院蘇真真醫師、饒敦醫師)的寶貴指點，謝謝大家。

中文摘要



自體免疫腦炎與副腫瘤腦炎是不常見但對患者有重要影響的疾病。其診斷之黃金標準為腦脊髓液或血清之抗體檢測，但檢查費時且價格並不便宜。因此，目前已有學者開發數種臨床量表評估抗體陽性機率，如 APE2 score 等等，惟這些量表均沒有腦波資訊。若能透過常規腦電圖 (electroencephalography, EEG) 篩選抗體陽性個案，預期可加速臨床決策。本前導初步研究以回溯性病例為基礎，收集 2017 年至 2022 年間臺大醫院疑似自體免疫腦炎病人 46 例（其中 12 例抗體陽性），評估單次靜息態 EEG 結合機器學習預測抗體陽性的可行性。

整體而言，本研究顯示陰性結果，發現以腦電圖單一模態之機器學習分類效能並不統計顯著優於隨機猜測。這顯示以自體免疫腦炎之單一腦電圖，以常見的特徵擷取與傳統機器學習分析方式的框架下，可能沒有明顯可以泛化(generalize)之特異發現。本研究結果有助釐清腦波在自體免疫腦炎診斷流程中的定位，可能為輔助性而非主要診斷性的角色，並提供後續腦波於免疫腦炎的研究比較基準。

關鍵字：腦波、機器學習、自體免疫腦炎、副腫瘤腦炎、自體免疫腦炎抗體

ABSTRACT



Autoimmune encephalitis and paraneoplastic encephalitis are uncommon yet clinically significant diseases. The diagnostic gold-standard is cerebrospinal-fluid autoantibody or serum antibody testing, a procedure that is time-consuming and costly. Several scoring systems—such as the APE2 score—have been proposed to estimate antibody positivity, but they do not incorporate electroencephalography (EEG) findings. If routine EEG could help identify antibody-positive patients, clinical decision-making might be accelerated.

46 patients who were admitted to National Taiwan University Hospital between 2017 and 2022 for suspected autoimmune encephalitis (AE). Among them, 12 were antibody-positive, making the prevalence 28%. The aim is to examine whether single resting state EEG, analyzed with machine-learning classifiers, could predict antibody status.

Overall, this pilot study demonstrates that machine learning models trained with EEG-only data do not perform better than random guessing significantly. These findings may hint that using common feature extraction methods and conventional machine learning algorithms, there may be no apparently generalizable discriminative features in one-time routine EEG. These negative findings suggest that the role of EEG in patients with suspected antibody-related encephalitis may not be directly diagnostic, but supportive. This study also provides a reference baseline for future research on EEG and immune encephalitis.

Keywords: electroencephalography, machine learning, autoimmune encephalitis, paraneoplastic encephalitis, autoimmune encephalitis antibodies

目次



誌謝	i
中文摘要	ii
ABSTRACT	iii
目次	iv
圖次	vi
表次	vii
Chapter 1 Introduction.....	1
1.1 自體抗體腦炎	1
1.2 自體抗體腦炎臨床預測分數	2
1.3 神經電訊號及腦波簡介(14).....	5
1.4 腦波於自體抗體腦炎之角色，暨研究目的	9
Chapter 2 Materials and Methods.....	11
2.1 病患來源及檢查項目	11
2.2 腦波前處理方法	12
2.2.1 Automagic automatic pipeline	12
2.2.2 Minimalistic pre-processing	15
2.3 腦波特徵擷取方法	16
2.3.1 傳統腦波特徵擷取方法	16
2.3.2 WEASEL-MUSE.....	19
2.4 機器學習框架	21

2.4.1	Binary classification supervised machine learning problem	22
2.4.2	Learning and Regularization.....	24
2.4.3	Cross validation and nested cross validation	25
2.5	機器學習方法	27
2.5.1	L2-regularized logistic regression (LR).....	28
2.5.2	Random forest (RF)	29
2.5.3	Support vector machine (SVM).....	31
2.5.4	Gradient boosting decision tree (GBDT).....	34
2.5.5	機器學習特殊參數選取與超參數調整	35
2.6	敘述統計.....	36
2.7	統計檢定力分析	36
Chapter 3	Results.....	37
3.1	病患特性.....	37
3.2	模型訓練結果	38
3.2.1	Automagic preprocessing	38
3.2.2	Minimalistic preprocessing.....	43
Chapter 4	Discussion	49
4.1	結論.....	49
4.2	研究限制.....	50
4.3	未來展望.....	50
	參考文獻	52

圖次



Figure 1.3-1: 國際 10-20 腦波電極位置系統(節錄自 (20)).....	7
Figure 1.4-1: 典型 extreme delta-brush (EDB)之腦波波型及頻譜圖(23).....	10
Figure 2.2-1: PREP 濾除市電前後之頻譜(29).....	12
Figure 2.3-1: Daubechies wavelet db-4 (37).....	19
Figure 2.3-2: Symbolic Fourier Approximation (SFA) (39).....	20
Figure 2.3-3: SFA 轉換對波型之影響(38).....	20
Figure 2.3-4: 一個維度的 WEASEL-MUSE 分析(38).....	21
Figure 2.4-1: ROC 圖例，取自 (41).....	23
Figure 2.4-2: 使用外部驗證挑選模型或參數(改自(40)).....	25
Figure 2.4-3 傳統交叉驗證(42).....	26
Figure 2.4-4: 巢狀交叉驗證(nested cross validation) (43).....	27
Figure 2.5-1: Soft-margin SVM 範例(取自(46)，圖中 β 為本文之 w).....	33
Figure 2.5-2: XGBoost 架構(取自(49)).....	35
Figure 3.2-1: 使用 Automagic 前處理的模型 balanced accuracy.....	39
Figure 3.2-2: 使用 Automagic 前處理的模型 AUC.....	39
Figure 3.2-3: 使用 Automagic 前處理的模型 F1 score.....	40
Figure 3.2-4: 使用最小化前處理的模型 balanced accuracy.....	44
Figure 3.2-5: 使用最小化前處理的模型 AUC.....	44
Figure 3.2-6: 使用最小化前處理的模型 F1 score.....	45

表次



Table 1.2-1: APE score (modified from (12)).....	3
Table 1.2-2: ONES score (modified from (12), part 1c)	3
Table 1.3-1: 突觸後電位與表深層對頭皮電位之影響(21)	7
Table 2.4-1: 常見二元預測的指標	22
Table 2.5-1: 本實驗測試之超參數範圍	35
Table 3.1-1: 抗體陽性與陰性病患特性	37
Table 3.1-2: 抗體陽性組之抗體細節	37
Table 3.2-1: 使用 Automagic 前處理的模型結果一覽表	40
Table 3.2-2: 使用 Automagic 前處理的模型 balanced accuracy 及統計比較	41
Table 3.2-3: 使用 Automagic 前處理的模型 AUC 及統計比較	42
Table 3.2-4: 使用最小化前處理的模型結果一覽表	45
Table 3.2-5: 使用最小化前處理的模型 balanced accuracy 及統計比較	46
Table 3.2-6: 使用最小化前處理的模型 AUC 及統計比較	47

Chapter 1 Introduction



1.1 自體抗體腦炎

在所有腦炎中，自體抗體陽性之腦炎(包含自體免疫腦炎(autoimmune encephalitis) (1)與副腫瘤抗體腦炎(paraneoplastic encephalitis) (2))，是相對少見(發生率介於十萬分之一至萬分之一人年(3)(4))但對病患影響重大的疾病。不同於感染性腦炎，他們是免疫系統攻擊腦部所產生的發炎性疾病。因此，在這群患者不應檢驗出明確致病之細菌、病毒、真菌，而且一般而言腦脊髓液之白血球數及蛋白濃度也不會如感染高(3)。自體抗體腦炎臨床上常表現為急性或亞急性發生的認知功能異常、局部神經症狀、或是癲癇(1)(2)，而確診的主要依據是在病患之腦脊髓液或血清中檢驗到特定抗體(5)。自體免疫腦炎主要是有對抗神經細胞表面之抗體(neuronal surface antigen)，副腫瘤抗體腦炎則是有對抗神經細胞內之抗體(neuronal intracellular antigen)，並且常與腫瘤相關，因此可能在診斷腦炎前後發現相對應的特異腫瘤(6)。致病機轉上，後者可能可以在腫瘤上找到原本在神經上表現的抗原，造成免疫系統誤判，稱作 molecular mimicry (7)。

過去經驗上，根據 Saraya 的報導，泰國醫療中心七年內(2010-2017)共收錄 77 位臨床診斷為腦炎的病患，其中 40%有單一表面(自體免疫腦炎)抗體，而 33%有單一胞內(副腫瘤腦炎)抗體(8)。兩者主要的差異在於前者發病年紀中位數較小，且行為表現的比例較高，但是其餘變項如性別、癲癇比例、腦脊髓液體分析、出院時失能狀況並沒有明確差異(8)。

因此，一位病患常只會檢驗到一個抗體，但是整體而言胞內與胞外抗體均有許



多種類，每一種抗體的表現都略有不同，較有特異性的臨床症狀或是檢驗結果也常被作為診斷參考(如下表 1-2.2)。確診之後，快速而正確的開始免疫治療是讓病患能夠恢復的重要因素之一(9)。然而，無特異臨床及檢驗(核磁共振、腦波等)表現之疑似自體抗體腦炎患者就必須等待抗體檢驗，如果抗體陰性，則要考慮進行腦切片生檢(biopsy)確診(10)，而這段診斷不明確的時間，是否先經驗使用免疫治療常令醫師及患者家屬為難。因此，檢驗抗體至結果出來之前，若是能夠事先評估抗體陽性機率高低，可望讓醫師更有信心，也更能說服家屬及早進行免疫治療，一邊等待檢驗結果，以期改善患者預後。

1.2 自體抗體腦炎臨床預測分數

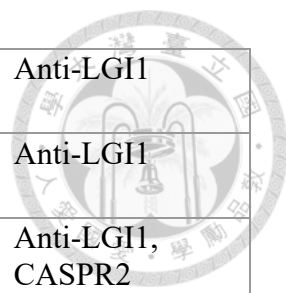
如上所述，因為疾病初期的症狀可以不特異(non-specific)，而且抗體檢驗耗時且並不便宜，所以有不少研究試圖使用臨床與腦部核磁共振的資訊做出評分系統，以提供抗體陽性率的先驗預測，如“APES (Antibody Presence in Epilepsy before Surgery)” (11), “APE2 (antibody prevalence in epilepsy and encephalopathy) score” (12) (下表 1.2-1), “ONES (obvious indications for neural antibody testing in epilepsy or seizures) score” (下表 1.2-2) (13) 等等。然而，以上評分系統雖有列入影像(如核磁共振的發現)，但是都沒有列入腦波(僅有癲癇發作(seizure)的特性，並無腦波，如下附表 1.2-1 及 1.2-2 所示)。另外，不同抗體之間較有特異性的臨床表現，幾乎完整被整理在表 1.2-2 中，以供臨床判斷。可以注意到，查無明顯原因之新發生癲癇發作，以及不明原因的癲癇惡化，在兩個量表搭配其他發現，就可能拿到足夠分數可以疑似自體抗體腦炎，所以本研究之主要研究族群為此類病患。

Table 1.2-1: APE score (modified from (12))

Criterion	Score
Autonomic dysfunction: 自律神經異常之發現	1
Brain MRI: 邊緣性腦炎的核磁共振發現	2
Seizure or cognitive changes: 一年內新發生之癲癇，或是六周內新發生且有進行性的認知功能變化	1
CSF findings 腦脊髓液發現合乎發炎性反應: 蛋白超過 50 mg/dl，白血球數量超過 5/dl，但紅血球需小於 1000 /dL	2
Facial dyskinesia 臉部不自主運動，尤其指臉手肌張不全性癲癇發作 (Faciobrachial dystonic seizure)這一種特殊癲癇發作	2
Malignancy 除卻單純皮膚的鱗狀上皮癌或基底細胞癌	2
Psychiatric symptoms 精神症狀	1
Seizure refractory to medical treatment 藥物難治之癲癇	2
Viral prodrome 發燒、喉嚨痛、流鼻水等病毒症狀，但因與某些癌症症狀類似，需無癌症才能計分	2

Table 1.2-2: ONES score (modified from (12), part 1c)

診斷條件	對應抗體
腦部核磁共振	
無萎縮之內側顳葉 FLAIR 高訊號病灶，與癲癇起始有時序相關性	Various
血管周邊之 Linear radial 顯影，與癲癇起始有時序相關性	Anti-GFAP, may overlap with anti-NMDAR
生化檢驗	
一年內新發生之顳葉或疑似顳葉癲癇，合併不明原因低血鈉 (定義為<130 mEq/L)	Anti-LGI1
臨床發現	
癲癇發作後不明原因之明確中樞神經或周邊神經病變	Various
音樂性癲癇	Anti-GAD65
臉手肌張不全性癲癇發作(Faciobrachial dystonic seizures)	Anti-LGI1



一年內新發生之難治型顳葉或疑似顳葉癲癇，併有毛髮直立性發作	Anti-LGI1
一年內新發生之難治型顳葉或疑似顳葉癲癇，同時併有陣發性頭暈	Anti-LGI1
一年內新發生之難治型顳葉或疑似顳葉癲癇，但是在五十歲之後才發病	Anti-LGI1, CASPR2
難治型顳葉或疑似顳葉癲癇，併有 anti-GAD65 相關之自體免疫疾病	Anti-GAD65
病史	
腫瘤診斷兩年內，新發生之癲癇	Various
使用腫瘤免疫治療(“immune checkpoint inhibitor”)一年內，心產生之癲癇	Various
單純疱疹病毒治療結束三個月內，新產生或惡化的癲癇	Anti-NMDAR



1.3 神經電訊號及腦波簡介(14)

因為神經元會維持細胞內外離子濃度之差異，細胞膜為脂質(lipid bi-layer)不讓離子任意流經，且物理特性接近電容器，所以神經靜止時內外之電位不同，這個差異為靜止膜電位(resting membrane potential)。神經元亦有多種不同平時關閉之離子通道，有些會與相合的化學物質結合而開啟(metabotropic ion channel)，有些會偵測細胞膜電位變化而開啟(ionophoric ion channel)，而這些離子通道開啟時會產生穿過細胞膜的電流，進而改變細胞膜的電位。因為鈉離子及鉀離子通道之特性，與胞內外不同之分布，神經可以被興奮而產生能夠向外傳遞的動作電位(action potential)。

神經元本身除了藉由細胞膜上的動作電位傳遞訊息，一個神經元也會在神經與神經交界處，使用突觸(synapse)藉由神經傳導物質(neurotransmitter)與突觸後神經元上的代謝型受體(metabotropic ion channel)結合產生突觸後神經電位(post-synaptic potential, PSP; 興奮性者稱為 excitatory PSP, EPSP; 抑制性者稱為 inhibitory PSP, IPSP)，來向突觸後神經元傳遞訊息。這些穿膜的離子電流也會導致細胞外電位的差異，所以吾人可以在細胞外觀察電位變化，來逆向推估神經的活動變化。

腦波是利用在頭皮表面黏貼電極，來非侵入性的紀錄頭皮表面電位。然而，大腦的皮質神經元與電極之間相隔甚遠，中間亦有腦膜、腦脊髓液、顱骨、頭皮等阻抗，所以量測到的訊號相當微小，而且若非有足夠多的神經元，有同時並持續足夠一段時間的訊號，表面電位的腦波是無法記錄到的。因此，一般認為，腦波之訊號產生並非來自時間短而不同步的單一神經動作電位，而是來自於持續時間久且位置比較大的突觸後電位所產生(15)。另外此一突觸電位源頭在深淺層對表面電位的影響相反(16)，而且興奮與抑制性突觸電位對表面電位的影響亦相反(整理於表 1.3-1)。在頭皮上，一般臨床腦波會有 19 個頭部電極，外加雙耳則為 21 個電極，根據

電極鄰近的腦部構造以及數字來命名，數字越小越靠中間，左側為奇數右側為偶數，如下圖 1.3-1。



由於癲癇發作是肇因於過度的腦部同步活動，而且睡眠腦波與清醒時之腦波樣式明確不同，所以腦波傳統是使用於確診癲癇並且監測追蹤用藥後患者的腦部活動狀態，以及睡眠多項生理檢查(polysomnography)等等。現今，除卻癲癇之外，腦波之應用範圍愈來愈寬廣，包含但不限於：

1. 認知行為研究與神經退化疾病(17) (18)

神經退化性疾病，如阿茲海默失智症與巴金森氏症等等，是人口老化之後的顯學。此時，是否能使用非侵入的腦波就得知神經退化疾病機轉相關的變化就相當重要。在(17)當中，作者試圖找尋不同種類退化疾病的腦波中，power law exponent β 是否會有差異。在(18)當中，作者比較各種不同臨床失智症診斷的患者之腦波差異。

2. 麻醉深度監測(19)

由於使用麻醉藥物之後病患的意識與反應會降低甚或消失，了解病患之麻醉深度不能藉由詢問病患得知，而必須仰賴客觀工具。因為過深的麻醉可能造成患者身體負擔，過輕的麻醉則可能造成病患不適，所以利用腦波於入睡與清醒時之不同以及特定藥物造成不同表現的特點，可以監測麻醉深度，進而能投與最佳劑量之麻醉藥。

3. 中風後監測(19)

有較大範圍缺血性或出血性之中風患者，因為有可能在中風後數日內產生腦水腫、缺血或出血範圍擴大、腦壓增大等問題而需要手術，所以病患往往會住在加護病房密切監測。然而，傳統監測方式為每小時將病人叫醒進行全



然而，腦波並非沒有限制(16):

1. 由於神經電訊號需要透過層層解剖結構傳出，每一層都有其等效電阻、等效電容等，所以腦波沒有辦法清楚提供深處腦部之資訊。因此，若要特別腦區之訊號，就必須使用具有侵入性的植入性電極，但是植入性電極則不能紀錄大範圍腦區之訊號，所以互有優缺點。
2. 腦波由於是紀錄頭皮表層訊號，所以容易受到各式生理與非生理因素干擾，前者包含頭臉部肌肉收縮的肌電、心臟來源之心電干擾、眼睛移動之眼動干擾、甚或流汗等等，而後者包含市電干擾、電極產生之干擾等等。後續章節會討論腦波不同前處理之方式來去除這些干擾。

3. 與腦部功能性磁共振造影(functional MRI, fMRI)比較(22):

腦波因為從腦部有電氣訊號至腦波變化幾乎沒有延遲，所以時間解析度只受儀器 sampling frequency 與放大器、濾波器的限制或選擇，但是空間解析度如上所述有所限制。fMRI 則相反，空間解析度可以自行選擇(FOV, Δx 等等)，但是因為 fMRI 是紀錄腦部血液含氧程度的 susceptibility difference，從腦部活動，血流變化(neurovascular coupling)，一路到訊號變化則有高達數秒之時間差，所以時間解析度相當有限。

4. 與腦磁圖(magnetoencephalography, MEG)的比較(22):

由於以上簡述的離子電流，其實這些微小的電流也會產生微小的磁場，其方向與電場方向垂直，所以也可以偵測磁場變化來推估腦部活動，此即為腦磁圖。腦磁圖的好處在於上述頭皮顱骨等生體解剖構造對磁場幾無影響，所以不容易受到扭曲。然而，因為腦磁相當微小，一般而論落在 femtotesla 的等級，所以需要非常靈敏的偵測器才能感應到這樣的變化，是故傳統腦

磁圖機置於 faraday cage 中，並使用超導體偵測。臨床上只有相當穩定且配合的病患才能施行，而且外來干擾的去除就相當重要；另一方面，肌肉收縮等干擾仍然會造成磁場訊號，所以並不會因為改用腦磁圖就消失。

1.4 腦波於自體抗體腦炎之角色，暨研究目的

目前公認具特異性的腦炎腦波發現僅有 NMDAR (N-methyl-D-aspartate receptor)腦炎所對應的 extreme delta-brush (EDB)。其主要表現為疊加在 1-3 Hz 規則之 delta 徐波(rhythmic delta activities)上的快速波(beta (20–30 Hz) 頻段波型) (23) (24)，典型之腦波與頻譜圖如下圖 1.4-1 (摘錄自(23))，可見在緩慢的背景波(藍色箭頭)上有明確的快速波型(紅色箭頭)。然而，其餘腦炎抗體至今並沒有具定論之敏感或特異之腦波特性(25)。

因此，探索腦波資訊是否能夠額外有助於預測腦炎抗體存在，在臨床診斷與治療上具有重要價值。考慮以上 extreme delta-brush 之發現，頻率分析方面的特徵萃取相當重要。本前導研究主要比較並使用傳統(傅立葉分析、小波分析、Hjorth parameters)，以及 symbolic Fourier transform (SFA)-based 之特徵擷取方式(feature extraction)，搭配不同機器學習演算法，試驗並探索是否可以使用腦波預測抗體存在的可行性。

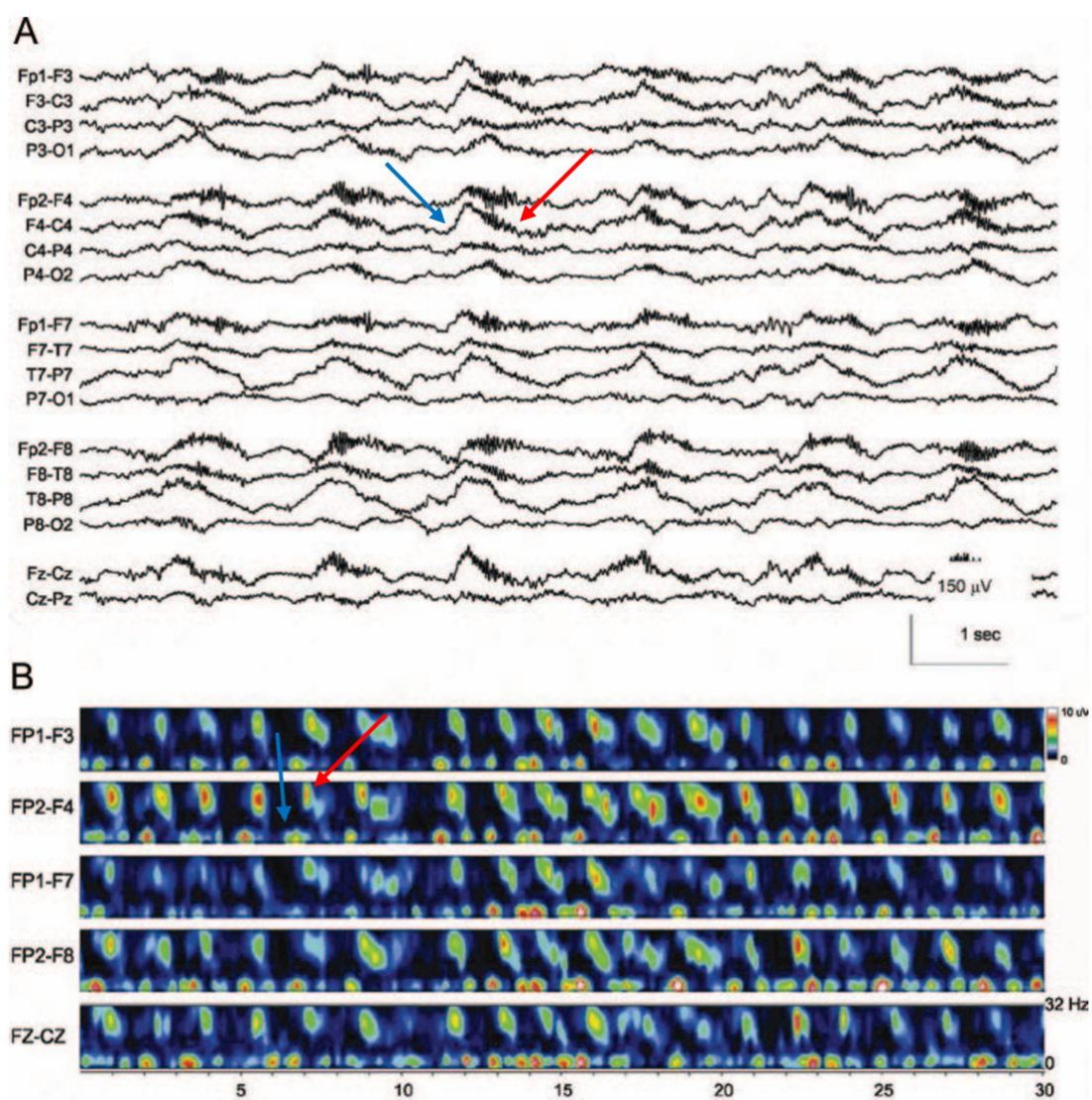


Figure 1.4-1: 典型 extreme delta-brush (EDB)之腦波波型及頻譜圖(23)

Chapter 2 Materials and Methods



2.1 病患來源及檢查項目

本研究回溯性收錄臺大醫院總院(國立臺灣大學醫學院附設醫院)於 2017 至 2022 年此五年內，因新發生癲癇發作(seizure)，或者不明原因癲癇惡化(unexplained aggravation of existing epilepsy)，懷疑可能為自體抗體腦炎而住院進行詳細檢查之病患，排除中樞神經感染、腦部腫瘤、近期中風、以及先前已知中樞神經發炎性之疾患，一共收錄 46 位病患，其離住院時間最接近的臨床腦波紀錄，並僅留取腦波中半段約五分鐘，已經檢查完電極確定沒有太多雜訊，但尚未進行過度換氣(hyperventilation, HV)或照光測試(intermittent photic stimulation, IPS)的段落。此段腦波接近一般之靜息態(resting state)腦波，但是病患可以依照自己的想法睜眼或閉眼。本實驗收錄之所有腦波均是由 Nihon Kohden 公司生產之 EEG-1200 系統腦波機，使用國際標準 10-20 電極位置紀錄(26)。

自體抗體檢驗係於華宇公司(Uni Pharma Co., Ltd)進行，自體免疫腦炎抗體項目包含 NMDAR (N-methyl-D-aspartate receptor)、CASPR2 (Contactin-associated protein-like 2)、LGI1 (Leucine-Rich Glioma-Inactivated Protein 1)、AMPA1/2 (α -amino-3-hydroxy-5-methyl-4-isoxazolepropionic acid receptor)、GABA/BR (γ -Aminobutyric acid receptor)六個項目，而副腫瘤抗體腦炎則包含 Hu、Ri、Yo、Tr、Ma2、amphiphysin、CV2、Recoverin、SOX1、Titin、Zic4、GAD-65 十二個項目。

以上回溯性收案內容業已經過國立台灣大學醫學院倫理中心倫理委員會之同意，案號為 202006184RINB。



2.2 腦波前處理方法

由於腦炎患者腦波常因意識不清、躁動導致腦波上的干擾，本研究中亦比較不同的前處理方式是否會對結果造成影響。

2.2.1 Automagic automatic pipeline

由於要一口氣處理大量腦波前處理並不容易，瑞士 Andreas 團隊提出 Automagic (27) 軟體，其整合數個常見在腦波處理軟體 EEGLab (28) 中的前處理 packages，包含 PREP pipeline (29)、ICLabel (30)，以及眼動圖回歸等等。

PREP pipeline 的初步處理，包含移除 60 Hz 市電干擾(line noise removal，在此 (29)先進行 sliding window (預設寬度為 4 秒，每個 window 差一秒鐘)，之後再於頻域擬合特定 60Hz 以及其諧波之 sinusoid component，之後去除該成分，濾波前後之頻譜圖如下圖 2.2-1 所示)、內部 re-referencing (流程中會先計算全體訊號平均，之後所有訊號減去平均)、訊號不良電極排除(包含震幅太大或太小(deviation criterion)、腦波電極訊號兩兩之間關係 (correlation criterion and predictability criterion，前者指兩兩電極應有一定的相關性，後者指兩兩腦波電極訊號應該有一定的互相預測能力)、雜訊(noisiness criterion)之條件，換算出一個 z-score，並用此判斷是否剔除不良電極)，後續再使用內插法補上排除的電極。

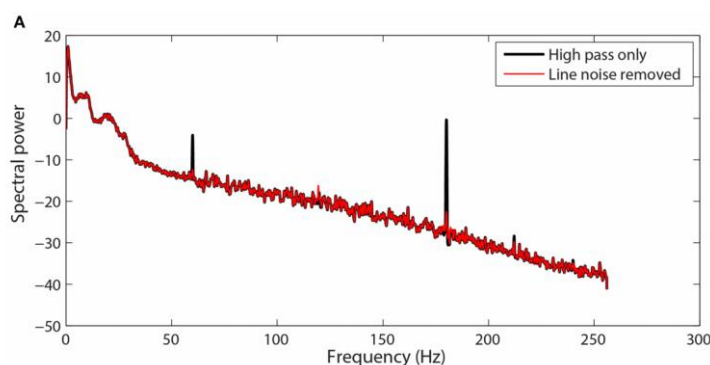


Figure 2.2-1: PREP 濾除市電前後之頻譜(29)

ICLabel (30)則是先使用獨立成分分析(31) (independent Component Analysis, ICA)將腦波拆解成數個獨立成分(independent component)，再使用事先訓練好的深度學習模型預測將預測符合眼動、肌肉收縮、心臟、電極雜訊或其他非腦部訊號的成分挑除。原始 ICLabel 之訓練資料集為超過二十萬份經由專家標註的獨立訊號，並已在過去實驗中顯示出可以有效除去雜訊。

獨立成分分析其實可以視作一種 matrix factorization algorithm，原理如下：假設訊號為矩陣 X (m -by- n matrix, m 個電極或是特徵、 n 個時間點(time points))，是由獨立的 k 個成分來源(寫成矩陣 S (k -by- n))的線性組合而產生：

$$X_{(m*n)} = A_{(m*k)}S_{(k*n)}, \text{ 或是 } S = WX, W = A^{-1}$$

由於中央極限定理，混合類似 resampling，所以後的矩陣 X 有較高機率會比原本假設之獨立成分矩陣更近似高斯分布(gaussian distribution)，因此可以設計藉由減少高斯分布的演算法以達到取回混合矩陣 A 與獨立成分 S 。著名 ICA 演算法其中之一 FastICA 的步驟如下(31)：

1. 先將原始矩陣 X 盡可能白雜音化(whitening)
 - a. 將原矩陣減去每一列(row)的平均，取得矩陣 X'
 - b. 對 X' 做奇異值分解(singular value decomposition), $X' = UDV^T$
 - c. 計算 $Z = D^{-1}U^T X'$ ，算出之矩陣 Z 的共變異矩陣就會變成 identity matrix，換言之 $ZZ^T = I$ 。
2. 針對 Z 進行計算，換言之找到 $S = WZ$ ，使得 S 的非高斯性(non-gaussianity)最高。方法有很多，比如說增加機率分布的 kurtosis 等。此處，FastICA 則是借助 negentropy 來達到此一目標。Negentropy 的想法是來自於高斯分布的熵(entropy)最大，所以盡可能遠離高斯分布熵的其他機率分布就會越不接



近高斯分布。由於對於合適的非線性函數 G ，一個隨機變數 y 的 negentropy J 會有以下特性(越大就越不接近高斯分布):

$$J \propto (E[G(y)] - E[G(v)])^2$$

FastICA 定義 $G(x) = \log \cos(x)$ ，並使用牛頓法找尋 $J' = 0$ 的解。所以會使用到 $g(x) = \tanh(x)$ 以及 $g'(x) = 1 - \tanh^2(x)$ 。假設我們先計算第一個獨立成分 s 所對應到的 W 內的係數 (column vector) w ，亦即 $s = w^T X$ ，那麼 J 對 w 微分，牛頓法則告訴我們遞迴計算下方的 w 可以找到 $J' = 0$ 的解，也就是原始 J 的極端值(31)：

$$w = w - \frac{J'}{J''} \approx E[XG'(w^T X)] - E[G''(w^T X)]w$$

以下演算法會逐次計算 W 每行(column)的數值，虛擬碼(pseudocode)如下，假定迴圈最多次數為 $\max_iter$ 。

- a. For p in 1 to k
 - i. Let new w = random column vector of length m
 - ii. $w = \frac{w}{\|w\|^2}$
 - iii. For count in 1 to $\max_iter$
 1. $w = \frac{1}{N} (Zg(w^T Z)^T - g'(w^T Z)1_N w)$, where 1_N is a column vector full of 1 of length N .
 2. $w = w - \sum_{i=1}^{p-1} (w^T w_i) w_i$, where $w_i = W[i, :]^T$
 3. $w = \frac{w}{\|w\|^2}$
 - iv. $W[p, :] = w^T$
- b. Return with W and $S = WX$.

眼動圖回歸則是將原始訊號對眼動圖做線性回歸，去除回歸的部分，以移除眼

睛傳來之干擾。做完上述前處理之後所取得之訊號，會再經過 0.1 Hz 高通濾波以及 average re-referencing (以所有電極訊號之平均做為參考電極，詳見下一段落)，再進行下一部份的機器學習。



2.2.2 Minimalistic pre-processing

為了瞭解進行複雜的前處理是否會造成以資判斷的訊號資訊流失或者扭曲，本實驗加入第二個前處理流程作為 automagic 的對照組，只保留最基本的雜訊移除，以比較不同處理策略對下一步機器學習預測能力的影響。細節步驟如下：

1. 市電干擾移除(60 Hz line noise removal)

本流程中只使用 EEGLAB 中最基礎的 cleanLine 函數去除此干擾。

2. 高通濾波器(0.1 Hz high pass filtering)

經過上述去除市電干擾的訊號，再使用 0.1 Hz 高通濾波器除去極低頻流汗或緩慢移動所造成的干擾。本實驗使用 EEGLAB 中的 pop_eegfiltnew 函數，使用 FIR 濾波器，最小 phase change 進行濾波。

3. 平均參考電極(Average referencing)

最後輸出之前，計算所有電極訊號之平均，做為參考電極，再進行下一部份的機器學習。

本方法並不會剔除不良通道(noisy channel)，也不會移除困難移除的雜訊，最大程度保留原始腦波訊息。另外，因為此流程的項目最少，設計初始想法也包含測試機器學習方法是否在雜訊中仍有預測抗體的能力，以及探索 real time processing 的可能性。



2.3 腦波特徵擷取方法

由於臨床腦波中，腦部有 21 個電極（不計入兩個眼動以及一個接地電極），檢測時間為十至十五分鐘，採樣頻率(sampling frequency)為 200 Hz，即便只取中間五分鐘的資料，資料數量仍然相當龐大(一位病患的腦波資料矩陣大小為 $(21, 6 \cdot 10^4)$)，除非使用其他機器學習已經訓練好的模型(pre-trained model)，傳統機器學習必須先進行特徵擷取(feature extraction)，才能避免參數過多的問題。本實驗中，以上至前處理部分是在 MATLAB 與 EEGLAB 架構下進行，而特徵擷取與後續機器學習分析則是在 python 進行。

2.3.1 傳統腦波特徵擷取方法

傳統特徵擷取方法包含以下三項，對每個電極都運算一次，使用 python 中之 numpy (32)、scipy (33)、以及 pywt(34)函式庫來實現。本實驗參酌(35)的特徵擷取方式設計。

1. Hjorth parameters (36)

Hjorth 提出三個分析腦波活性的方式，假設其中一個電極的資料為 $y(t)$ ：

a. Activity = $\text{Var}(y(t))$

定義就是此電極訊號的變異數。

b. Hjorth Mobility = $\sqrt{\frac{\text{Var}\left(\frac{dy(t)}{dt}\right)}{\text{Var}(y(t))}}$

此數字會與 power spectral standard deviation 有關，所以可以使用時域數字預估訊號頻域的狀況。

c. Hjorth Complexity = $\frac{\text{Mobility}\left(\frac{dy(t)}{dt}\right)}{\text{Mobility}(y(t))}$



如果訊號是完全的弦波(pure sine wave)，這個數字會是 1，若訊號與弦波差異越大，與此數字就會離 1 越遠。

因此，一個電極會算出三個特徵數字。

2. DFT spectrum

考慮連續週期性的訊號，傅立葉分析之基礎奠基於 $\left\{\frac{e^{inx}}{\sqrt{2\pi}}\right\} \forall n \in \mathbb{Z}, x \in [-\pi, \pi]$ ， $(e^{ix} = \cos(x) + i \cdot \sin(x))$ 可以形成正交基底(orthonormal basis)，所以任意訊號 $f(t), t \in [-\pi, \pi]$ 都可以投影在這組基底上以找出其傅立葉係數(Fourier series coefficients)：

$$F(n) = \langle f(t), \frac{e^{int}}{\sqrt{2\pi}} \rangle = \frac{1}{\sqrt{2\pi}} \int_{-\pi}^{\pi} f(t) e^{-int} dt$$

而此轉換的最大優點為 sin 與 cos 是形狀簡單的週期性訊號，所以其轉換後的虛數長度有直覺的解釋，即「原始訊號在該頻率的強度」。離散傅立葉轉換則是以上對週期性訊號分析的離散版本。

在本實驗中，此處會計算一個電極在所有時間範圍內的離散傅立葉轉換，但只取其長度，並不取相位(phase component)。假設訊號為 f，k 為 DFT 之後的 index，假設訊號長度為 n：

$$\text{Pre-Feature}_k = \text{abs} \left\{ \sum_{m=0}^{n-1} f(m) \exp \left(-2\pi i \frac{mk}{n} \right) \right\}$$

根據(35)的特徵擷取方式，本實驗中使用 numpy 中的 FFT，但並不是直接計算出結果就直接列為特徵，而是在以上頻譜的 pre-features 中挑選出最大值、最小值、平均、標準差、變異數、第 5,25,50,75,95 個百分位數、以及均方根，以上總共十一個特徵。為了避免以上方式會遺漏目前已知 delta-beta 之間的關係，本實驗中亦直接加入 delta band (定義為 0 ~ 4 Hz) 與 beta band (定義

為 13 ~ 30 Hz)相對應的 pre-feature 線下積分值(亦即 area under amplitude spectrum curve)。以上總共有十三個特徵。以上是將整段五分鐘的腦波直接做傅立葉分析，沒有做額外的 windowing，由於取樣頻率為 200 Hz，頻率解析度為 0.0033 Hz per sample。

3. Wavelet analysis

本實驗中使用多貝西小波(Daubechies wavelet)中的 db-4 來進行小波分析，延續過去研究的說法(37)，如圖 2.3-1。原始連續小波分析的公式如下：

$$\text{CWT}(\tau, s) = \frac{1}{\sqrt{|s|}} \int f(t) \psi^* \left(\frac{t - \tau}{s} \right) dt$$

小波轉換是時頻轉換，有時間點 τ 以及 scale s 。本實驗中則是使用其離散版本的離散小波轉換，概念與連續相當類似，但是不同於後者使用的是 scales，前者主要是將訊號通過事先設計好(符合可以重現原訊號、正交等條件，切分頻率落於取樣頻率一半)的高通與低通濾波器，再降採樣(downsampling)以得到細節的高頻與粗略的低頻成分，這樣是一個 level。對於低頻成分，可以再遞迴進行進一步離散小波分析，level 指的就是遞迴的次數。本實驗使用五個 levels，再將這五個 levels 的計算值連接一起。同樣並非以直接轉換後的數字做為特徵，而是在以上結果中挑出最大值、最小值、平均、標準差、變異數、第 5,25,50,75,95 個百分位數、以及均方根，全部共十一個特徵。以上亦是將整段五分鐘腦波直接納入分析，並沒有做額外的 windowing。

因此，在傳統特徵擷取方式中，一個電極會產生 27 個特徵(一份腦波總共會產生 $21 \cdot 27 = 567$ 個特徵)以供後續機器學習探索。

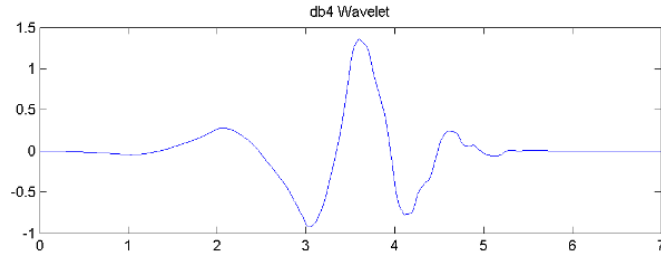


Figure 2.3-1: Daubechies wavelet db-4 (37)

2.3.2 WEASEL-MUSE

WEASEL-MUSE (38) 是 “*Word ExtrAction for time SEries cLassification plus Multivariate Unsupervised Symbols and dErivative*” 的縮寫，是針對多個時間序列 (multiple time series，例如腦波，一份腦波有多個電極資料) 的處理方法。原理上，它是從 symbolic Fourier approximation (SFA) (39) 所演衍伸出來，SFA 的目的在於將時間序列化做文字進行後續分析：

1. 給定一個時間序列 $\{y(t)\}_{t=0}^{n-1}$ ，對它做離散傅立葉轉換，可以得到與原序列長度一樣的傅立葉係數 $\{Y(k)\}_{k=0}^{n-1}$ ，對係數取實部與取虛部後共有 $2n$ 個係數。
2. 對同性質 p 個時間序列都進行上述分析，就可以得到 $2np$ 個數字，接下來就可以做 multiple coefficient binning (MCB)。以 equi-depth binning 為例，觀察這 $2np$ 個數字在 $2n$ 行(column) 中的狀況，在有一定數值個數時插入分隔(breakpoint)，就可以將原本的連續的係數離散化(discretization)成文字 (SFA word)，如下圖 2.3-2。這個作法是大略可逆的，由文字可以取回傅立葉係數在 binning 中的中點，由這些係數做反轉換就可以得到近似的序列，如圖 2.3-3，但其解析度明顯會受到 binning 的影響。

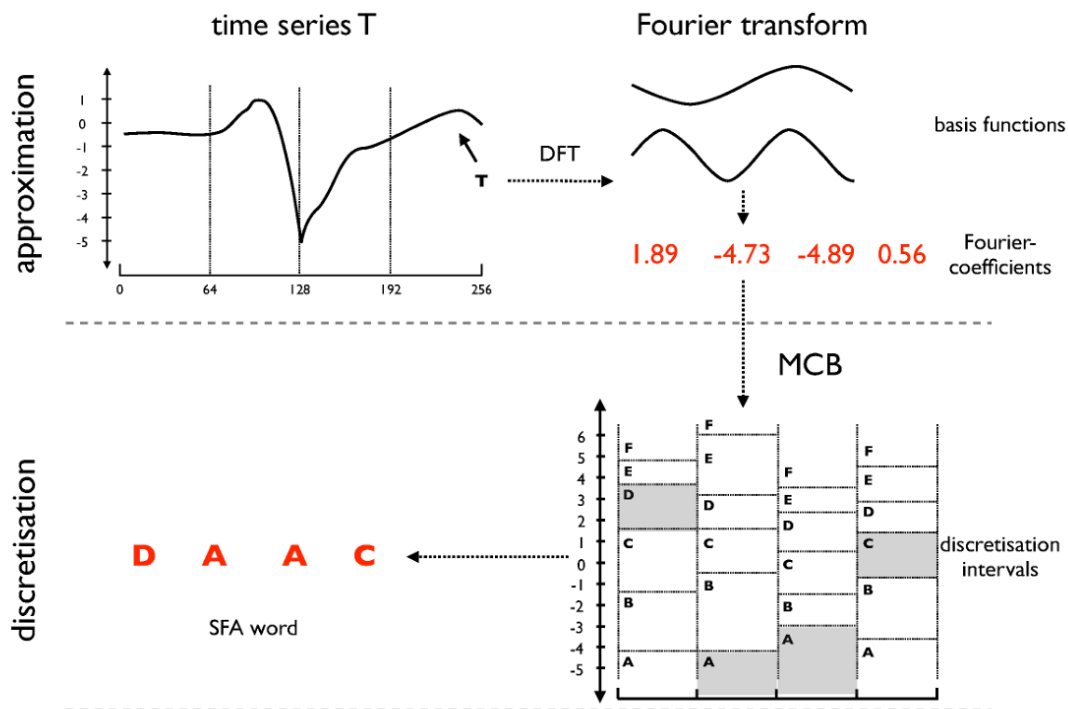


Figure 2.3-2: Symbolic Fourier Approximation (SFA) (39)

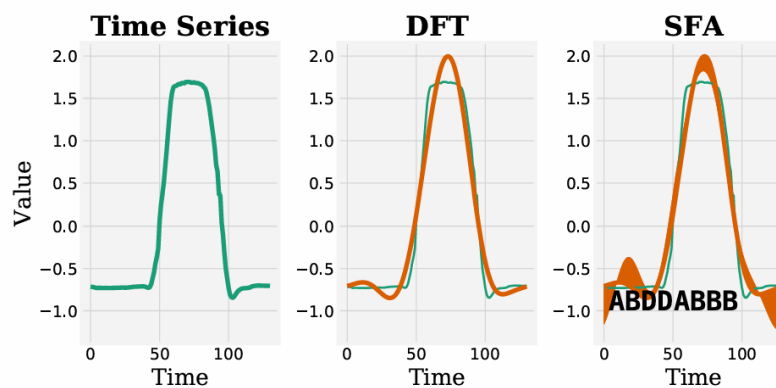


Figure 2.3-3: SFA 轉換對波型之影響(38)

WEASEL-MUSE 則是對 SFA 的延伸，針對多個時間序列，以腦波為例資料維度為(病患數，電極數，時間點數)，它會逐一對給定的維度(在此顯為電極那一個維度)下去分析。但每一個維度之分析方式與先前 SFA 稍有不同，它會先將資料切出 windowing，再對 windowed data 做傳統 SFA，而轉換成的文字最後再產生 histogram，如圖 2.3-4。此演算法會最後會對每一個維度各產生一個 bag-of-pattern histogram。

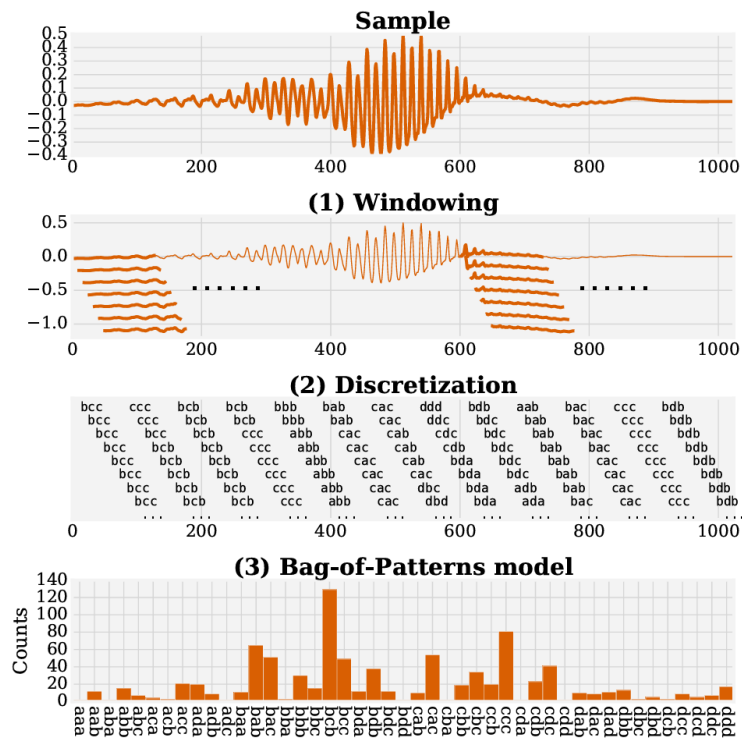


Figure 2.3-4: 一個維度的 WEASEL-MUSE 分析(38)

最後產生的每一個維度之 bag-of-pattern histogram，將資料維度調整為(病人數，電極數*bin 數目)，就可以直接進入機器學習流程。本實驗中嘗試除預設值(預設值為 4)外的 bin size 與 word size 與 window_sizes，確保 beta 頻段都有收錄(window 使用原資料長度的 0.1, 0.3, 0.5, 0.7, 0.9 倍，每個倍數各跑一次，沒有重疊(window overlap)；最大嘗試之 word size 為 2000，對應 1000 個傅立葉虛數，在最小 window 時最大可以涵蓋至 33 Hz；bin size 最大嘗試 16)，但是對機器學習結果都沒有影響。

2.4 機器學習框架

本實驗因是預測抗體陽性或陰性，所以是二元預測監督式機器學習問題(binary classification supervised machine learning problem)，整體架構如下(40)：

2.4.1 Binary classification supervised machine learning problem

假設訓練資料集(training dataset)是 $\{(X_i, y_i)\}_{i=1}^N$ ，其中有 N 筆資料， X_i 為第 i 個樣本的特徵，而 y_i 為第 i 個樣本的二元預測目標， $y_i \in \{-1, 1\}$ 代表陰性或陽性，我們假設背後有個函數 f 使得 $y_i = f(X_i)$ 。任一機器學習模型之目的在於學習結果 g 能夠逼近 f ，使得此模型可以用以預測未來新而未知的資料。因此，若將機器學習演算法比擬為考生，二元預測資料集比擬為是非題考題，整體架構概念上約略同於先讓考生看過去的考古題，要求考生分析如何判斷答案，之後再拿全新的考題讓考生作答，也可以此驗證是否考生真的有學到判斷的依據。

一般而言，吾人是使用傳統的統計分析方式來判斷二元預測系統的好壞，如同考卷的評分方式。方法包含 accuracy, balanced accuracy, sensitivity (Recall), specificity, positive predicted value (PPV; precision), negative predicted value (NPV) 等等，如下表 2.4-1。如果資料有偏頗(意即陽陰比例離一半相當遠)，就不適合使用單純的 accuracy，因為模型直接猜 majority class 就可以拿到高 accuracy，這時需要使用 balanced accuracy，或是 $F1\ score = \frac{2}{recall^{-1} + precision^{-1}}$ 才可以避免這個問題。至於 AUC score 需要模型支援輸出預測的機率才能運算，每一個資料點對應的 class probability，按照該點真實為陽性或陰性可以推成該點的 sensitivity and specificity，畫在縱軸為 sensitivity，橫軸為 1-specificity (意即 FPR) 的 ROC (receiver operating characteristic) 圖上，就可以擬和曲線來計算線下面積，如下圖 2.4-1。

Table 2.4-1: 常見二元預測的指標

	預測陽性	預測陰性	
--	------	------	--

真陽性	True positive (TP)	False negative (FN)	Sensitivity (Recall, Sen, TPR) $= \frac{TP}{TP + FN}$
真陰性	False positive (FP)	True negative (TN)	Specificity (Spe) $= \frac{TN}{TN + FP}$
Accuracy (全部總和 N) $= \frac{TP + TN}{N}$	PPV (Precision) $= \frac{TP}{TP + FP}$	NPV $= \frac{TN}{TN + FN}$	Balanced accuracy $= \frac{Sen + Spe}{2}$

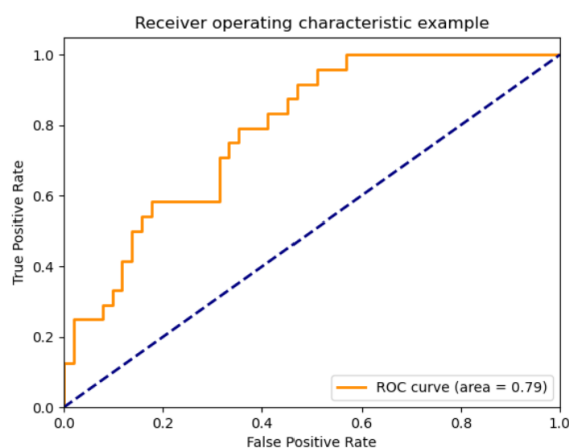


Figure 2.4-1: ROC 圖例，取自 (41)

然而，機器學習模型內部卻未必會使用以上的指標來引導訓練，而是事先定義一個包含比較模型輸出與正確答案之間差異(例如 mean squared error, MSE)，以及模型特定(model-specific)項目的損失函數(loss function)。以上目的在於達到輸出值與正確答案之間差異越小越好，但是同時可以達到模型正則化(regularization)之目的，後者之理由詳見下一段落。



2.4.2 Learning and Regularization

根據 Vapnik-Chervonenkis 的理論，我們可以使用 VC dimension 的來度量機器學習的模型空間 $\mathcal{H}$ 的複雜度。此一維度可以假想成模型的參數(effective parameters)數量；如果模型越複雜，參數越多，則此值越大。

假設機器學習模型是在 VC dimension= d_{vc} 的模型空間 $\mathcal{H}$ 中找到一個函數 g 來回答吾人長度為 N 的訓練資料集，那麼函數 g 在訓練資料集的錯誤率 $E_{in}(g)$ ，與函數 g 在真實世界分布中的錯誤率 $E_{out}(g)$ ，對任意 $\epsilon > 0$ ，有以下關係(VC bound; no free lunch theorem)(40)：

$$P(|E_{in}(g) - E_{out}(g)| > \epsilon) \leq 4(2N)^{d_{vc}} e^{-\frac{1}{8}\epsilon^2 N}$$

以上式子表示函數 g 無法泛化(generalize)的機率(訓練資料集錯誤率與真實錯誤率之間相差甚大的機率)，在 d_{vc} 非無窮大且 N 足夠大的時候，會被式子右側所限制住。換言之，如果 d_{vc} 有限， N 越大，則學出來的函數 g 在訓練資料集中的表現與真實世界中的表現就會趨同。因此，機器學習在統計上是站得住腳的。

然而，上式也說明一個狀況，如果訓練資料集大小 N 不夠大，若還是要有類似的泛化能力，就需要調降 d_{vc} 。在許多機器學習的實作中，就是使用正則化(regularization)以減少實質參數(effective parameters)的數量，來達到這個目的。

常見的正則化方式包含 L1 (LASSO) and L2 (weight decay)，若模型的參數為一個 column vector w ，前者是在損失函數中增加 $\alpha|w|$ ，後者則是增加 $\beta w^T w$ ，其中 α 與 β 是可調整的超參數(hyperparameter)，調整正則化效果之強弱，也是在作最佳化時的 Lagrange multiplier。因為 L1 有不可微分之處，所以計算較為麻煩，但是它因為邊上有 constant slope 所以容易產生稀疏解(sparse solution)，而 L2 則反之，通常較為容易計算但是較不會產生稀疏解。



2.4.3 Cross validation and nested cross validation

一般而言，如果有額外長度為 K 的外部驗證資料集做外部驗證(external validation)，使用模型在驗證資料集的錯誤率 $E_{\text{val}}(g)$ ，經過訓練後模型輸出的函數 g 在真實世界的表現就會很高機率會有(40)：

$$E_{\text{out}}(g) \leq E_{\text{val}}(g) + O\left(\frac{1}{\sqrt{k}}\right)$$

相同的，也可以使用外部驗證(external validation)做模型或參數挑選，假設有 M 個模型，其模型空間為 $\{\mathcal{H}_i\}_{i=1}^M$ ，在這個額外長度為 K 的外部驗證資料集 D_{val} 作驗證，挑最好的模型 g_{m^*} ，如下圖 2.4-2，那麼就很高機率有(40)：

$$E_{\text{out}}(g_{m^*}) \leq E_{\text{val}}(g_{m^*}) + O\left(\sqrt{\frac{\ln(M)}{K}}\right)$$

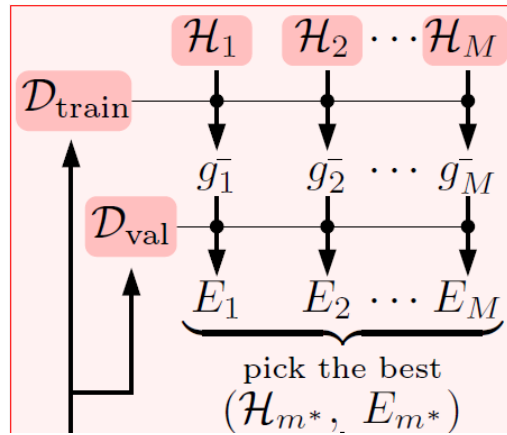


Figure 2.4-2: 使用外部驗證挑選模型或參數(改自(40))

因此，使用資料來驗證模型能力與調整參數是有理論依據的。然而，實務上，通常沒有那麼多資料集可以一個拿來訓練、另一個調整參數、最後再一個拿來做模型能力評估，大部分會使用交叉驗證(cross validation)。傳統交叉驗證是使用來調整參數或選擇模型，其想法在於完整使用全部的訓練資料集，每一個部分都曾作為訓練資料集與驗證資料集的一部分，但是最後拿來做模型能力評估的 test data 則完全

不可碰觸以避免洩題(data leakage)，如下圖 2.4-3。

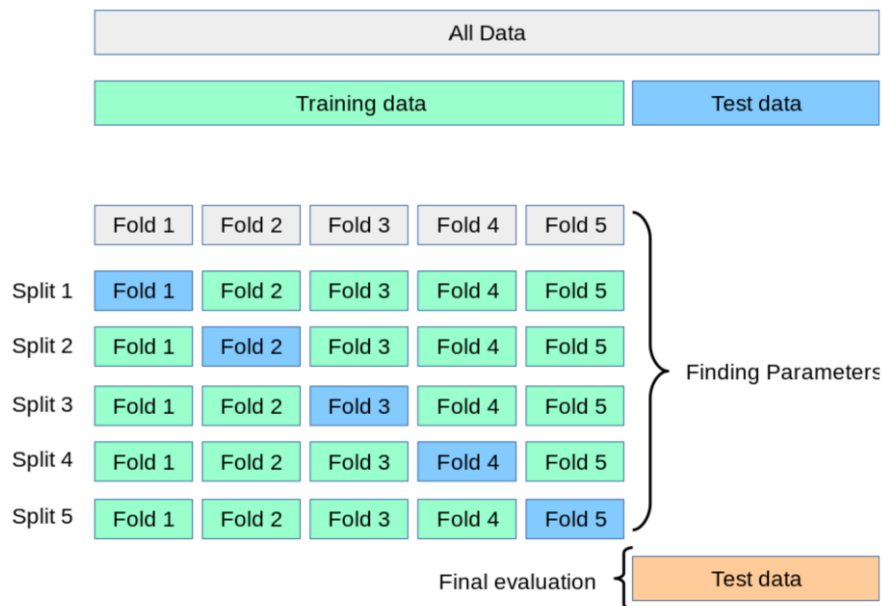


Figure 2.4-3 傳統交叉驗證(42)

因為本實驗的樣本數偏少(46)、腦炎患者並不容易收錄，又需要調整模型之超參數，所以不適合直接採用上述留出法(hold-out test data)來進行機器學習的成效分析(因為 hold-out test data 的個數會太少)，而只能採用巢狀交叉驗證(nested cross validation)，以外層的交叉驗證估算模型的正確率，並在內層的交叉驗證以 GridsearchCV 找尋最佳之超參數，最大程度避免 data leakage，但又能夠讓所有資料都盡可能物盡其用，既做過訓練資料，也做過調整參數的對照組，以及驗證正確率的測試組，如下圖 2.4-4。

另外，由於本實驗使用交叉驗證設計，而且腦波具有前後關聯性，並不合適搭太過簡易的 upsampling 方法，例如將五分鐘腦波切分成每一分鐘一個段落，而將樣本數量提增五倍。這樣個體自己的差異(intrapersonal variation)明顯比較小的狀況，

機器不需要學到真正的判斷方法，只要看過某個體的其中一個段落腦波就容易答對該個體其他段腦波的答案，造成假性分數膨脹，因此本實驗暫不考慮取此方法。

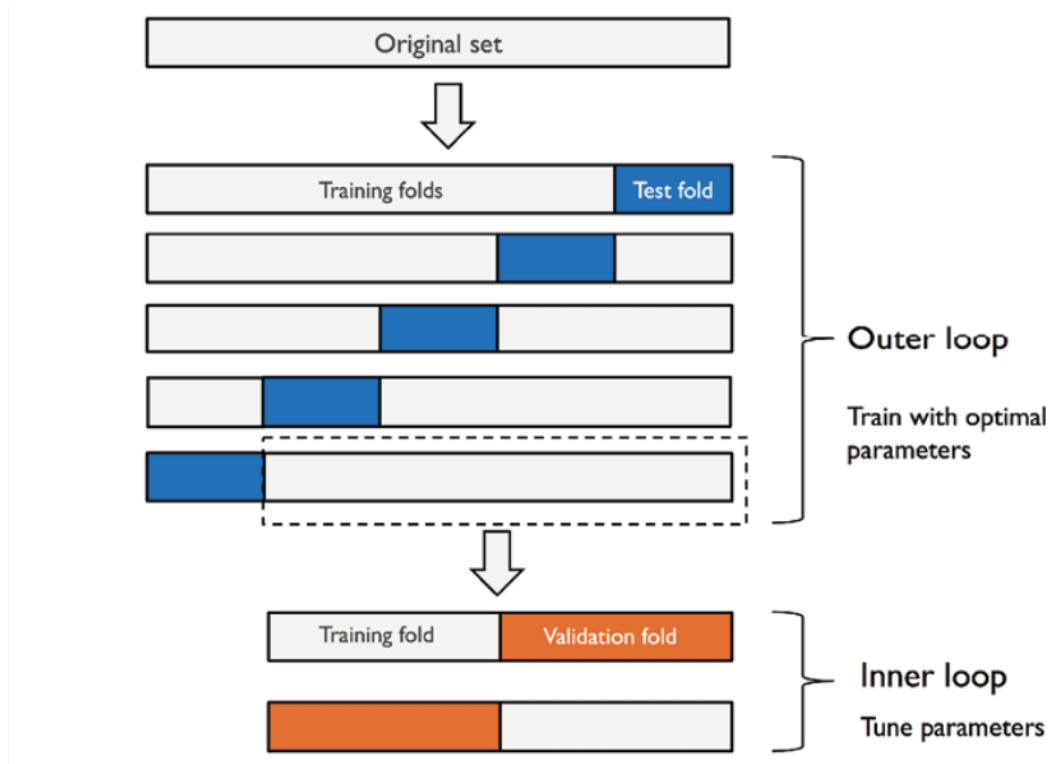


Figure 2.4-4: 巢狀交叉驗證(nested cross validation) (43)

在本實驗中，外圈的交叉驗證(outer loop)切做五份，而內圈的交叉驗證(inner loop)切做三份(並非上面示意圖中的兩份)，外圈用以評估模型預測能力，內圈挑選參數，而挑選參數時所使用之指標為 balanced accuracy。另外，在機器學習之前，訊號會先經過 standard scaler 減去平均除以標準差，以避免不同電極間的差距影響機器學習演算法。

2.5 機器學習方法

如上一節所述，因為樣本數偏少，而且目前並沒有大型針對腦波訓練之 pre-trained model，所以本實驗中並沒有採用現在時興的 deep neural network (DNN)的

模型訓練，如 InceptionTime (44)，而是使用傳統的機器學習方式，細節如下(40)(41)：

2.5.1 L2-regularized logistic regression (LR)

Logistic regression 其實是廣義線性模型(generalized linear model)的一部分。假設模型參數向量為 w ，一筆輸入資料為向量 x (此向量有插入常數 1 作為第一個值)，則模型輸出為 x 的線性組合結果套上一個 sigmoid function θ ：

$$\theta(w^T x) = \frac{1}{1 + e^{-w^T x}}, \text{ where } \theta(s) = \frac{1}{1 + e^{-s}}$$

此模型的損失函數則是 cross-entropy loss with L2 regularization，假設訓練資料集為 $\{(X_i, y_i)\}_{i=1}^N$ (在此所有 X_i 均插入常數 1 作為第一個值， $y_i \in \{0, 1\}$)，正則化參數為 C ，則有：

$$\text{Loss} = \frac{1}{C} \|w\|^2 - \sum_{i=1}^N y_i \log(\theta(w^T X_i)) + (1 - y_i) \log(1 - \theta(w^T X_i))$$

此模型之損失函數為 cross entropy 之原因在於 logistic regression 假設套過 sigmoid 函數的輸出項為 y 是陽性或是陰性的二項式分布機率(binomial probability)，且不同樣本之間互相獨立，所以似然率函數(likelihood function)為：

$$L = \prod_{i=1}^N \theta(w^T X_i)^{y_i} (1 - \theta(w^T X_i))^{(1-y_i)}$$

要使模型最貼切觀察數據，就要將似然率函數提高。由於 $\log$ 是單調遞增函數，取 $\log$ 並不會影響方向。由於傳統上損失函數希望越低越好，所以取負值，結果就產生：

$$\text{CE} = -\log(L) = -\sum_{i=1}^N y_i \log(\theta(w^T X_i)) + (1 - y_i) \log(1 - \theta(w^T X_i))$$



2.5.2 Random forest (RF)

隨機森林最早是由 Breiman (45)所提出，是決策樹(decision tree)加上 bagging with uniform blending 而產生。以下先從每次分支都固定 C 個子節點的一顆決策樹開始討論。一般而言，給定一個樣本 x ，分支條件函數 $b(x) \in \{1, \dots, C\}$ ， $[]$ 代表 indicator function (若輸入為真則輸出 1，反之則輸出 0)， $G_c(x)$ 代表第 c 個分支的子樹，則這棵樹的輸出決定函數可以寫成：

$$G(x) = \sum_{c=1}^C [b(x) = c] G_c(x)$$

其中，CART 決策樹是相當傳統的機器學習法，給定訓練資料集 $\{(X_i, y_i)\}_{i=1}^N$ ，其參數 $C=2$ ， b 則是 decision stump (shifted scaled Heavside function)，而每個分支的目的在於讓子樹的裡面的 y_i 越齊一 (impurity 越低) 越好，而定義 impurity 為 Gini index $= 1 - \sum_{k=1}^N \left(\frac{\sum_{n=1}^N [y_n=k]}{N} \right)^2$ ，訓練的虛擬碼如下(pseudocode)：

- Function CART ($D = \{(X_i, y_i)\}_{i=1}^N$):
 - If not branchable:
 - Return constant value = majority y_i label in D
 - Else:
 - Find $b(x) = \operatorname{argmin}_{h(x)} (\sum_{c=1}^2 |D(h, c)| \cdot \text{Impurity}(D(h, c)))$, where $D(h, c) = \{(X_n, y_n) \mid h(X_n) = c, (X_n, y_n) \in D\}$.
 - Split D into 2 parts:

$$\{D_c\}_{c=1}^2, \text{ where } D_c = \{(X_n, y_n) \mid b(X_n) = c, (X_n, y_n) \in D\}$$
 - For c in 1 to 2:
 - Build sub-tree with recurring call: $G_c = \text{CART}(D_c)$
 - Return $\sum_{c=1}^2 [b(x) = c] G_c(x)$

如果以上“branchable”的定義是沒有樣本而無法再向下分割，最後吾人得到的就是 full-grown decision tree。很明顯的，這樣的樹會使得訓練資料集的分類完全正確，換言之 $E_{in} = 0$ 。相反地，這樣複雜度高的模型通常不容易在真實資料上泛化 (generalize)，所以實務上會使用比較嚴苛的 branchable 條件使得分支提早結束。

隨機森林則是將原始訓練資料集，利用抽後放回的方式重新採樣(resampling with replacement)，來產生多組(以下以 B 組為例)略有不同的新資料集(稱作 bagging)，每個資料集都訓練一顆決策樹 $\{\phi_1, \phi_2, \dots, \phi_B\}$ ，最後預測時則讓所有的樹，一樹一票投票出最後預測結果(稱作 uniform blending)，換言之最後輸出的函數為：

$$g(x) = \frac{1}{B} \sum_{i=1}^B \phi_i(x)$$

隨機森林的原理若以回歸為例會比分類更容易理解(實際上，二元分類問題即是回歸 target 限縮成兩個數字)，在回歸問題中，使用 mean squared error 作為損失函數，假設未知的真實函數為 f ，所以對於訓練資料集 $D = \{(X_i, y_i)\}_{i=1}^N$ 有 $y_i = f(X_i) + \epsilon$ ，且 $\epsilon \perp X_i, E[\epsilon] = 0, \text{Var}[\epsilon] = \sigma^2$ ，而模型最後學到的函數為 g ，則根據 bias-variance decomposition (40) (46)：

$$E_{\text{out}} = E_D E_{X,y} [(y - g(X))^2] = \sigma^2 + E_X [E_D [(g - E_D[g])^2] + (f - E_D[g])^2]$$

其中第二項(中間項，也可以寫作 $\text{Var}_D[g]$)定義為 variance，和本次學到的函數與平均學到的函數之間的距離有關，第三項定義為 bias，和真實函數 f 與平均學到的函數有關。如果假設隨機森林中的每一棵樹有同樣的變異數 $\text{Var}_D[\phi_i] = \sigma_\phi^2$ ，而且樹兩兩之間有同樣的共變異數 $\forall i \neq j, \text{Cov}_D[\phi_i, \phi_j] = \rho$ ，那麼：

$$\text{Var}_D[g] = \text{Var}_D \left[\frac{1}{B} \sum_{i=1}^B \phi_i(x) \right] = \frac{1}{B^2} \left(\sum_{i=1}^B \text{Var}_D[\phi_i] + \sum_{i \neq j} \text{Cov}_D[\phi_i, \phi_j] \right)$$

$$= \frac{1}{B^2} (B\sigma_\phi^2 + B(B-1)\sigma_\phi^2\rho) = \sigma_\phi^2(\rho + \frac{1-\rho}{B})$$

上面的式子就會變成：

$$E_{\text{out}} = \sigma^2 + E_X[\sigma_\phi^2(\rho + \frac{1-\rho}{B}) + (f - E_D[g])^2]$$

在(47)中，Louppe 證明以上的 ρ 等於訓練資料集的變異數除以總變異數，而且 $\rho \geq 0$ ，所以隨機森林在 resampling 時如果越隨機， ρ 就接近零，隨著樹的數木 B 增多，假設其他數值變化不大， E_{out} 就會下降，模型的泛化能力就會逐漸顯現出來。

2.5.3 Support vector machine (SVM)

硬性邊界線性支援向量機(hard-margin linear SVM)的特色是在資料線性可分類(存在超平面可以將二元資料完全分對組)時，不是任意選擇一個可分對類的超平面(如最傳統的 perceptron learning algorithm)，而是選擇一個盡可能遠離資料點的超平面，由於新增這個限制，使模型複雜度下降，達到模型容易泛化的效果(40)。

換言之，考慮二元分類問題，假設訓練資料集為 $\{(X_i, y_i)\}_{i=1}^N, y_i \in \{-1, 1\}$ ，超平面為 $f(X) = w^T X + b$ ，線性支援向量機會在以下兩者之間取折衷(40) (46)：

1. 分類正確，對所有 i ， y_i 和 $w^T X_i + b$ 要同號。這點可以看出 (w, b) 的 scale 不重要，對於一組符合條件的 (w, b) ，任意乘上一個大於零的 k 的 (kw, kb) 也會符合條件，所以實務上這個條件可以改為 $\forall i \in \{1 \dots N\}, y_i(w^T X_i + b) \geq 1$ 。
2. 資料點盡可能離超平面越遠越好： $(w, b) = \arg \max_{(w, b)} \min_i \frac{|w^T X_i + b|}{|w|}$

因為 (w, b) 線性的出現於分子跟分母，可以再次發現 scale 不重要，所以可以不失一般性的額外假設 $\min_i |w^T X_i + b| = 1$ ，那麼上述式子就會變成 $w =$

$$\arg \max_w \frac{1}{|w|} = \arg \min_w w^T w。$$

3. 上述以上兩者放在一起就是：



$$\min_w w^T w \text{ s.t. } \forall i \in \{1 \dots N\}, y_i(w^T X_i + b) \geq 1$$

可以使用 quadratic programming 的方式來求解。也可以發現，只有落在

$y_i(w^T X_i + b) = 1$ 的資料點才會影響 (w, b) ，所以它們被稱作 support vector。

如果資料不是線性可分，但是原理上還是有接近線性的邊界(例如：因為雜訊的關係兩組資料出現重疊，但是原理上是線性邊界)，就可以使用軟性邊界線性支援向量機(soft-margin linear SVM)，這是上面硬性邊界的改良版，它將第一個條件放鬆為對所有 i ， $y_i(w^T X_i + b) \geq 1 - \xi_i$ ； $\xi_i \geq 0$ 。如此一來，如果資料點落在邊界上，(如下圖 2.5-1 的點 7)或是被分對類(如點 8)，則 $\xi_i = 0$ ，反之 $\xi_i > 0$ 。整體的模型是既要讓邊界最寬，但又盡可能分對類(盡可能減小 ξ_i)的折衷，中間有個超參數 C 可以調整兩者的權重(40) (46)：

$$\min_w w^T w + C \sum_{i=1}^N \xi_i \text{ s.t. } \forall i \in \{1 \dots N\}, y_i(w^T X_i + b) \geq 1 - \xi_i; \xi_i \geq 0$$

以上最佳化問題可以使用 quadratic programming 來解，也可以使用 Lagrange multiplier 將問題轉至 dual space 來求解。

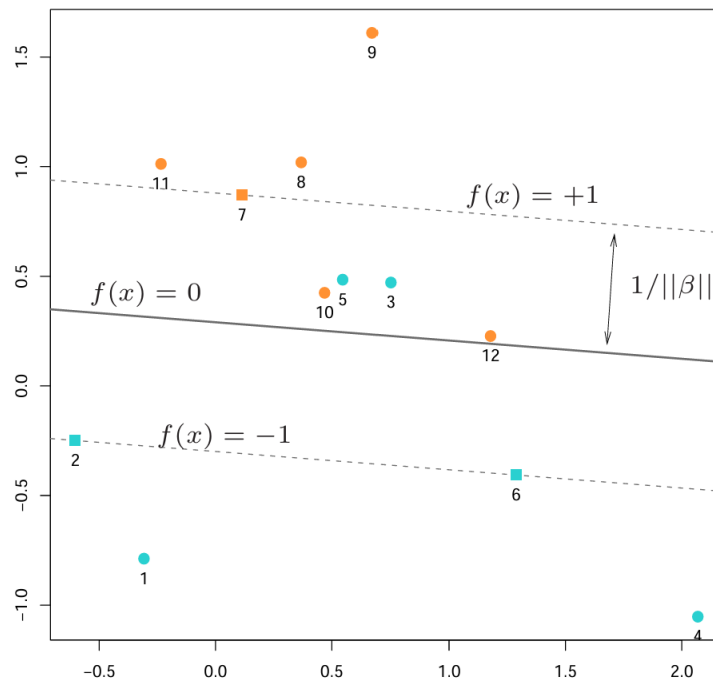


Figure 2.5-1: Soft-margin SVM 範例(取自(46)，圖中 β 為本文之 w)

如果資料非線性可分類的，就必須套用核函數(Kernel function)，先將原始資料轉換至高維空間內，或許在這個高維空間內資料是線性可分類的。對於任意兩個資料點 X_i, X_j ，核函數 K 有以下條件(Mercer theorem)，則核函數可以視作某種轉換函數 ϕ 的內積($K(X_i, X_j) = \phi(X_i)\phi(X_j)$) (40) (46)：

1. 對稱： $K(X_i, X_j) = K(X_j, X_i)$
2. 半正定(positive semidefinite)：給定任意 $\{c_i\}_{i=1}^N$ ，有 $\sum_{i=1}^N \sum_{j=1}^N c_i c_j K(X_i, X_j) \geq 0$

觀察以上就可以使用軟性邊界線性支援向量機的式子，因為它實際上 L_2 regularized linear model，所以使用 representer theorem 與 dual form 就可以套用核函數 kernel trick，最佳化以下公式就可以得到 kernel SVM (40) (46)：

$$\min_{\{\alpha\}} \sum_{i=1}^N \sum_{j=1}^N \alpha_i \alpha_j y_i y_j K(X_i, X_j) - \sum_{i=1}^N \alpha_i$$

$$\text{s. t. } \sum_{i=1}^N \alpha_i y_i = 0; \quad \forall i \in \{1, 2, \dots, N\} \quad 0 \leq \alpha_i \leq C$$



2.5.4 Gradient boosting decision tree (GBDT)

隨機森林是使用重新抽樣訓練資料集來產生不同的決策樹，所以每棵樹的產生不會有互相影響，GBDT 則是讓逐顆訓練新的決策樹，讓它學習前面已經生成好的所有樹的預測值與真實值之間的殘差(如下圖 2.5-2)，最後再讓訓練好的決策樹的結果依照設定好的權重 blending 輸出，此權重可以每顆樹相同，也可以跟樹的正確率有關(46)。

最常見的 XGBoost (48) 的原理如下，假設訓練資料集為 $\{(X_i, y_i)\}_{i=1}^N$ ，已經訓練好的 $t-1$ 顆決策樹總和的預測是 $\widehat{y_{t-1}}$ ，現在要訓練第 t 顆決策樹 f_t ，那麼就是要最小化以下的損失函數：(其中， $\text{err}(y_i, \hat{y}_t)$ 可以是上述 cross entropy error 或是 mean squared error，因為它是事先先被選好的，所以微分等數值都能計算； T_t 代表 g_t 總共有幾個子葉(leaf)， w_i 則是代表各個子葉的分數，定義在最後； γ 與 λ 是可以調整的超參數，使模型不容易過度複雜，達到正則化的效果)

$$L = \sum_{i=1}^N \text{err}(y_i, \widehat{y_t(X_i)}) + \Omega(f_t), \text{ where } \hat{y}_t = \widehat{y_{t-1}} + f_t; \Omega(f_t) = \gamma T_t + \frac{1}{2} \lambda \sum_{i=1}^{T_t} w_i^2$$

而使用到 Gradient descent 的部分是來自於泰勒展開式，以及選擇 w_j 使得損失函數能夠達到最小：

$$\text{err}(y_i, \widehat{y_t(X_i)}) \approx \text{err}(y_i, \widehat{y_{t-1}(X_i)}) + g_i f_t + \frac{1}{2} h_i f_t^2,$$

$$\text{where } g_i = \left. \frac{\partial \text{err}(y_i, x)}{\partial x} \right|_{x=\widehat{y_{t-1}}(X_i)} \text{ and } h_i = \left. \frac{\partial^2 \text{err}(y_i, x)}{\partial x^2} \right|_{x=\widehat{y_{t-1}}(X_i)}$$

$$\text{For any leaf } j \text{ and data inside it } D_j, \text{ set } w_j = -\frac{\sum_{i \in D_j} g_i}{\sum_{i \in D_j} h_i + \lambda} \text{ so that } \frac{\partial L}{\partial w_j} = 0$$

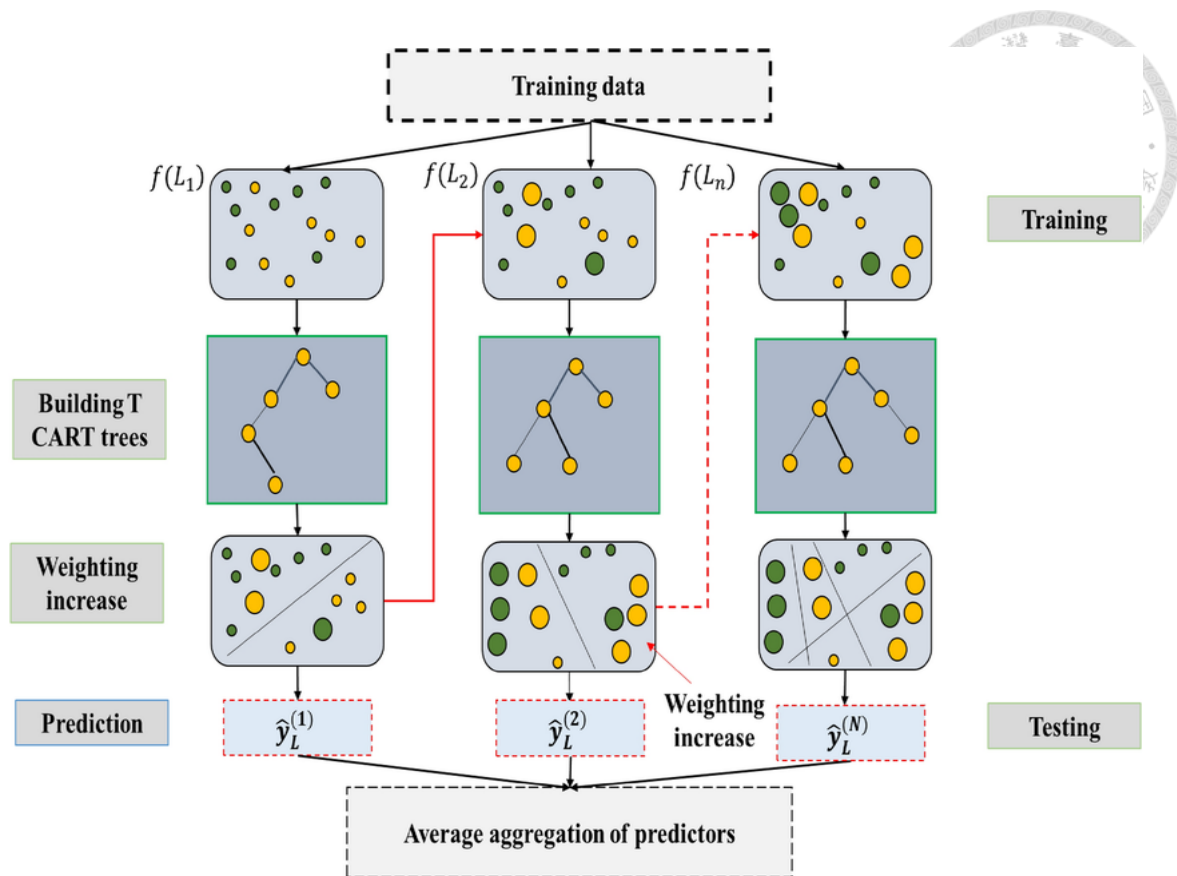


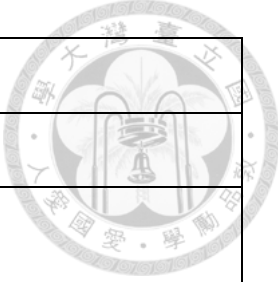
Figure 2.5-2: XGBoost 架構(取自(49))

2.5.5 機器學習特殊參數選取與超參數調整

本實驗使用 GridsearchCV 調整超參數，其搜尋範圍如下表 2.5-1。由於抗體陽性與陰性的比例不均，所以所有的交叉驗證都使用分層交叉驗證(stratified cross validation)以確保交叉驗證中沒有出現偏頗、所有機器學習演算法均有調整類別權重(class weight)以盡量避免機器只學到猜測抗體陰性組別、而且主要的評量指標為平衡正確率(balanced accuracy)。另外，由於本實驗樣本少但腦波特徵多，為了不要過擬和，所以沒有額外使用 SMOTE (50)等方式進行 oversampling。

Table 2.5-1: 本實驗測試之超參數範圍

Logistic regression	C: [1E-3, 1E-2, 1E-1, 1, 10, 100]
Random forest	N_estimator: [3, 9, 27, 81, 100]



	max_depth: [3, 5, None]
SVM (linear)	C: [1E-2, 1E-1, 1, 10, 100]
SVM (RBF)	C: [1E-2, 1E-1, 1, 10, 100] gamma: [1E-2, 1E-1, 1, 10]
XGBoost	max_depth: [3,6,9,12] min_child_weight: [1E-2, 1E-1, 1, 10]

2.6 敘述統計

除卻以上的腦波處理以及機器學習之外，本實驗亦簡單敘述統計統整此 46 位病患之特性，以及模型最後結果，類別變項是以 fisher exact test 來比較，連續變相先以 shapiro-wilk 檢驗常態性，資料若不能反駁常態性則使用 t-test 比較平均值，若不符合常態分佈則使用 Wilcoxon (Mann-Whitney U) test 進行無母數分析中位數。

2.7 統計檢定力分析

使用(51)的公式計算，以本研究收錄的病患陽性率 28%而言，假定預期第一型錯誤率(type I error) $\alpha = 0.05$ ，統計檢定力(power) $1 - \beta = 0.8$ ，以 0.5 作為虛無假設的 AUC，則機器學習的能力，以 AUC 評估必須要大於 0.75 左右，才能在本實驗人數觀察到顯著效果。

相對的，若 AUC 只有 0.7，那麼統計檢定力會衰退到 0.6 左右；若預期模型能力在 0.65 或 0.6，統計檢定力甚至會低至 0.37 或 0.19。換言之，本研究無法有效鑑別 AUC 低於 0.7 的機器學習成效。

Chapter 3 Results



3.1 病患特性

本研究一共收錄 46 名病患。其中，34 位抗體陰性，而 12 位抗體陽性，兩個群體的人數並不均等，約四分之三的病患均為抗體陰性。然而，兩個組別的男性比例、住院年紀、以及過去是否曾有癲癇診斷並沒有顯著差異，如下表 3.1-1。

Table 3.1-1: 抗體陽性與陰性病患特性

	抗體陰性組	抗體陽性組	p-value
人數	34 (74%)	12 (26%)	
男性比例	12 (35.3%)	6 (50.0%)	0.49
住院年紀	47.9 ± 19.0 (47.5)	47.2 ± 22.3 (49.5)	Wilcoxon: 0.78
過去有癲癇診斷	12 (35.3%)	3 (25.0%)	0.72

這十二位抗體陽性的病患中，最常見的抗體是 Anti-NMDAR 抗體，共有七位病患有此抗體，其餘的分布就相當繁雜，也有一位病患同時具有多個抗體，如下表 3.1-2。

Table 3.1-2: 抗體陽性組之抗體細節

Patient	Antibody / antibodies
1	Anti-NMDAR
2	Anti-NMDAR
3	Anti-NMDAR



4	Anti-NMDAR
5	Anti-NMDAR
6	Anti-NMDAR
7	Anti-NMDAR
8	Anti-CASPR2
9	Anti-GAD65, anti-titin, anti-recoverin, anti-CV2
10	Anti-Ma2/Ta
11	Anti-GABABR and anti-LGI1
12	Anti-GABABR

3.2 模型訓練結果

3.2.1 Automagic preprocessing

經過 Automagic 自動化前處理流程所產生的腦波特徵(Features)，在進行第二章所敘述之機器學習模型訓練與交叉驗證後，整體表現皆接近隨機猜測，無法有效有效預測抗體陽性與否，如下表 3.2-1、3.2-2、3.2-3，以及圖 3.2-1、3.2-2、3.2-3。

先論傳統搭配傳統腦波特徵擷取方法(包含 Hjorth 參數、傅立葉與小波頻譜)，以平衡正確率(balanced accuracy)與 AUC 而言，雖然最佳組合為線性支援向量機(linear support vector machine, linear SVM)，但是它的平均平衡正確率卻只有 0.595，AUC 亦約 0.605，顯示模型的鑑別能力僅略高於隨機猜測，而且統計分析也發現平衡正確率並不顯著優於隨機猜測(one sample t-test against 0.5)。其他機器學習方法，包含使用 RBF 核函數的支援向量機與 GBDT 等模型，平均平衡正確率與 AUC 大多均落在 0.5 至 0.6 之間。直接以傳統統計比較平衡正確率與 AUC，可以發現所有模型都無法顯著大於 0.5，即便其中有些模型的數值因為不符合常態分佈而必須

使用 Wilcoxon test 比較中位數，其中位數也仍然不顯著大於 0.5。各模型在 F1 score 分數的表現，最佳者約為 0.41，均不理想。

WEASEL-MUSE 之特徵擷取方法的整體模型表現更差，所有機器學習方法都在猜 majority class，所以平衡正確率均為 0.5，標準差為零，以至於 t-test 無法計算；其中若是 precision 或 recall 為零的模型，scikit learn 連 F1 score 都不能計算(實際上可以當作零)。這個結果反映出，在本資料集中，WEASEL-MUSE 比起傳統擷取方式擷取出的特徵更差，並不適合處理本實驗的類別不平衡小樣本資料集。

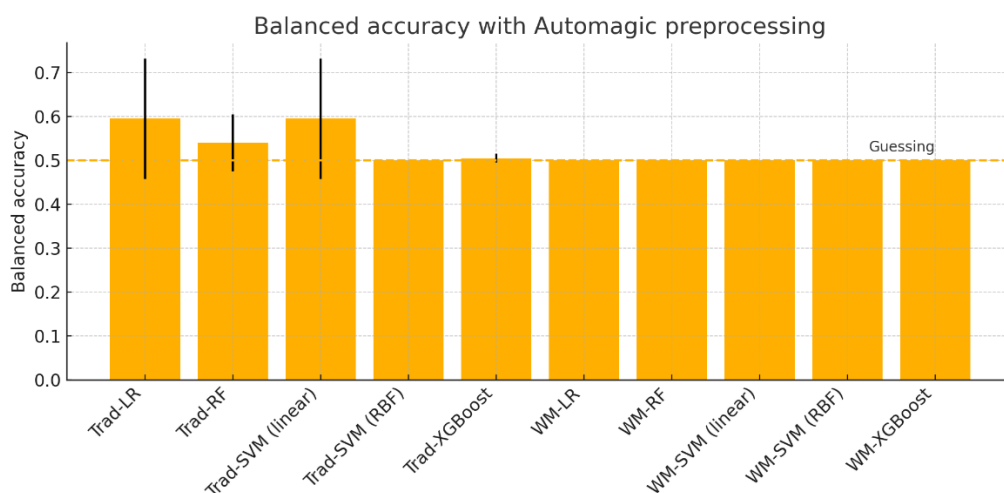


Figure 3.2-1: 使用 Automagic 前處理的模型 balanced accuracy

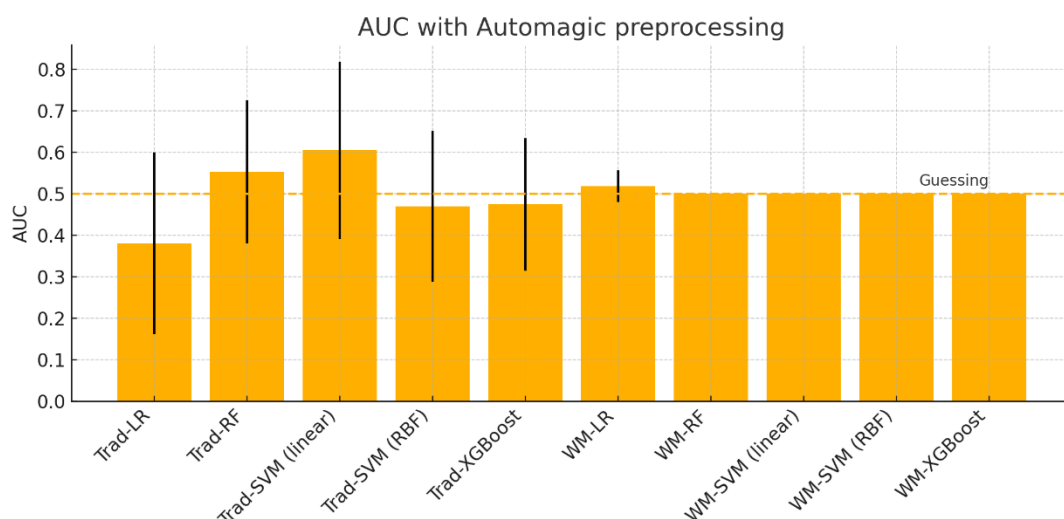


Figure 3.2-2: 使用 Automagic 前處理的模型 AUC

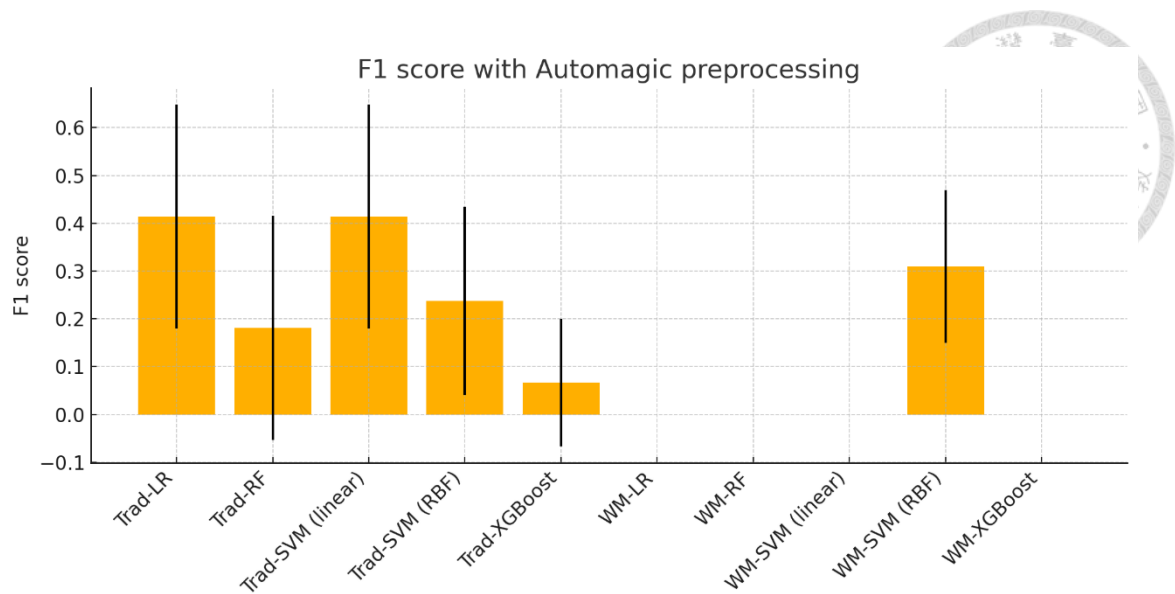
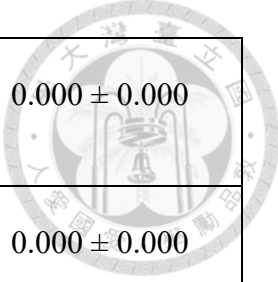


Figure 3.2-3: 使用 Automagic 前處理的模型 F1 score

Table 3.2-1: 使用 Automagic 前處理的模型結果一覽表

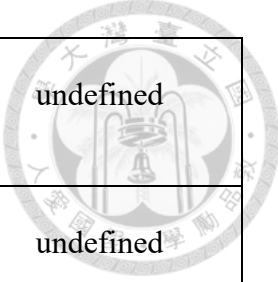
Automagic	Balanced accuracy	AUC	F1 score
Traditional Logistic regression	0.595 ± 0.137	0.381 ± 0.219	0.414 ± 0.234
Traditional Random forest	0.540 ± 0.065	0.553 ± 0.172	0.181 ± 0.234
Traditional SVM (linear)	0.595 ± 0.137	0.605 ± 0.214	0.414 ± 0.234
Traditional SVM (RBF)	0.500 ± 0.000	0.470 ± 0.181	0.238 ± 0.197
Traditional XGBoost	0.505 ± 0.010	0.475 ± 0.160	0.067 ± 0.133
WM Logistic regression	0.500 ± 0.000	0.519 ± 0.038	0.000 ± 0.000



WM Random forest	0.500 ± 0.000	0.500 ± 0.000	0.000 ± 0.000
WM SVM (linear)	0.500 ± 0.000	0.500 ± 0.000	0.000 ± 0.000
WM SVM (RBF)	0.500 ± 0.000	0.500 ± 0.000	0.310 ± 0.160
WM XGBoost	0.500 ± 0.000	0.500 ± 0.000	0.000 ± 0.000

Table 3.2-2: 使用 Automagic 前處理的模型 balanced accuracy 及統計比較

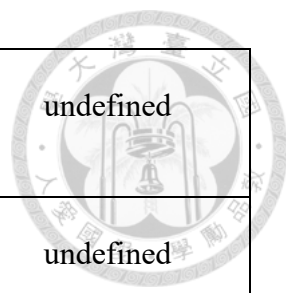
	Balanced accuracy	Shapiro-Wilk p-value	T-test (against 0.5) p-value
Traditional Logistic regression	0.595 ± 0.137	0.4543	0.2379
Traditional Random forest	0.540 ± 0.065	0.00635*	0.2784 (t-test) 0.1797 (Wilcoxon)
Traditional SVM (linear)	0.595 ± 0.137	0.4543	0.2379
Traditional SVM (RBF)	0.500 ± 0.000 (all data being 0.5)	1	undefined
Traditional XGBoost	0.505 ± 0.010	0.00013*	0.3739 (t-test) 0.3173 (Wilcoxon)
WM Logistic regression	0.500 ± 0.000 (all data being 0.5)	1	undefined



WM Random forest	0.500 ± 0.000 (all data are 0.5)	1	undefined
WM SVM (linear)	0.500 ± 0.000 (all data are 0.5)	1	undefined
WM SVM (RBF)	0.500 ± 0.000 (all data are 0.5)	1	undefined
WM XGBoost	0.500 ± 0.000 (all data are 0.5)	1	undefined

Table 3.2-3: 使用 Automagic 前處理的模型 AUC 及統計比較

	AUC	Shapiro-Wilk p-value	T-test (against 0.5) p-value
Traditional Logistic regression	0.381 ± 0.219	0.3920	0.3376
Traditional Random forest	0.553 ± 0.172	0.4905	0.5700
Traditional SVM (linear)	0.605 ± 0.214	0.2493	0.3837
Traditional SVM (RBF)	0.470 ± 0.181	0.2762	0.7561
Traditional XGBoost	0.475 ± 0.160	0.6527	0.7663
WM Logistic regression	0.519 ± 0.038	0.00013	0.3739 (t-test) 0.3173 (Wilcoxon)



WM Random forest	0.500 ± 0.000 (all data are 0.5)	1	undefined
WM SVM (linear)	0.500 ± 0.000 (all data are 0.5)	1	undefined
WM SVM (RBF)	0.500 ± 0.000 (all data are 0.5)	1	undefined
WM XGBoost	0.500 ± 0.000 (all data are 0.5)	1	undefined

3.2.2 Minimalistic preprocessing

經過最小化(Minimalistic preprocessing)流程所處理之腦波，整體而言，機器學習模型的成果大致上比 Automagic 更差(更近於隨機猜測)，如下表 3.2-4、3.2-5、3.2-6，以及下圖 3.2-4、3.2-5、3.2-6。

傳統腦波特徵擷取出的特徵，無論搭配哪一個機器學習方式，其平衡正確率與 AUC 幾乎全部相當接近 0.5，而且線性支援向量機在此的平衡正確率甚至統計顯著小於 0.5，劣於亂猜。WEASEL-MUSE 則與上述 automagic 的結果相當類似，模型都在預測 majority class，以至於平衡正確率都是 0.5，且 t-test 無法計算。F1 score 則有更多模型因為其 precision 或 recall 為零而軟體無法計算。雖然最小化前處理可以保留更多原始訊號成分，但很明顯的是本實驗的特徵擷取方式原理上無法區辨來源為腦部訊號、非腦部的生理訊號或非生理雜訊，所以本實驗發現，在雜訊移除程度不夠的前提下，模型更難以辨識判斷的規則。

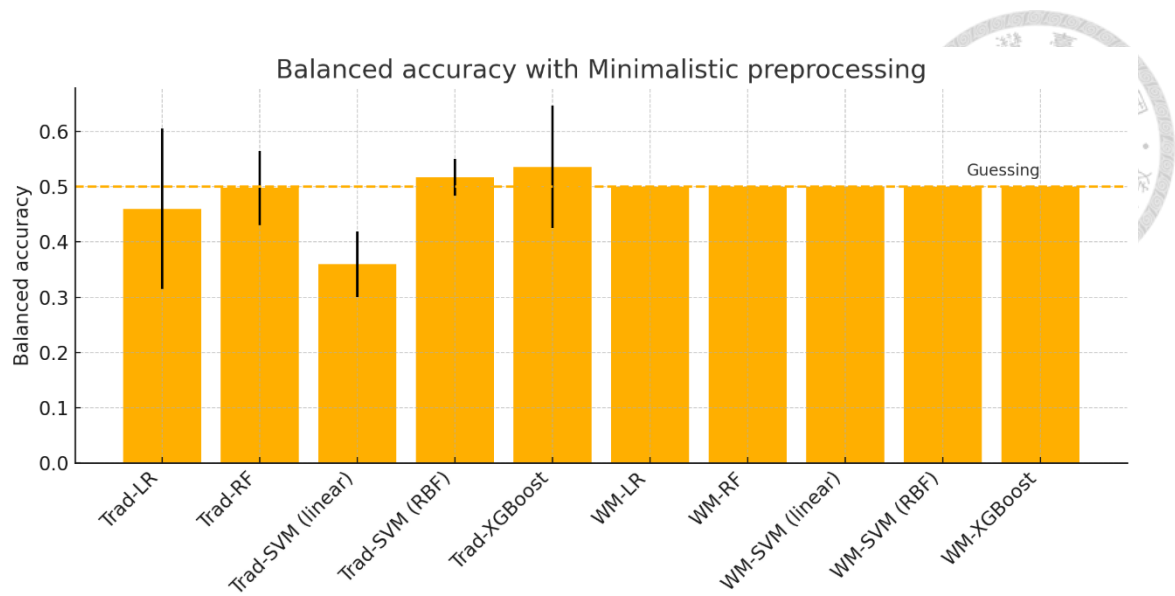


Figure 3.2-4: 使用最小化前處理的模型 balanced accuracy

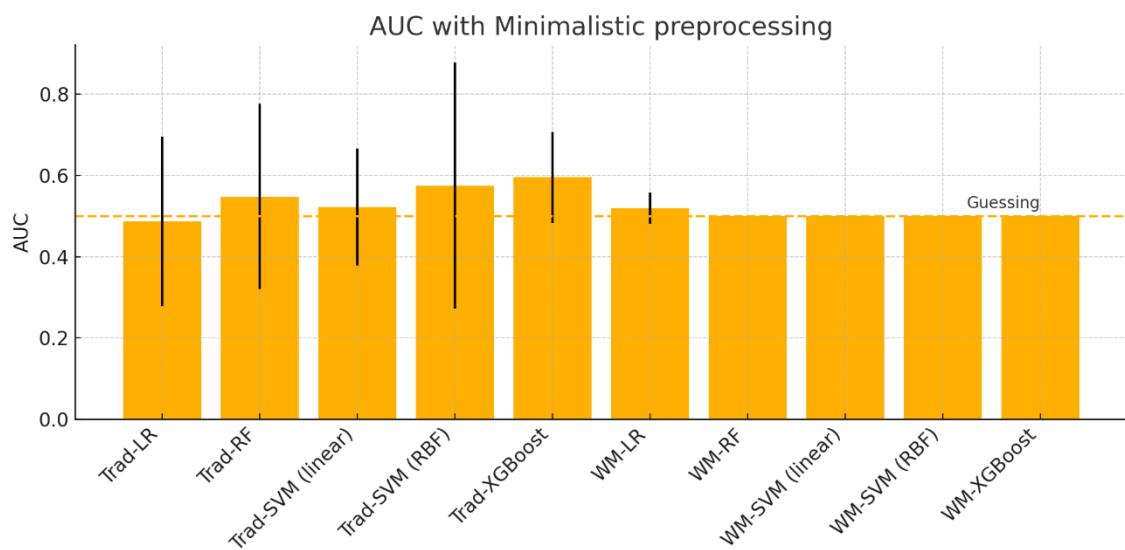


Figure 3.2-5: 使用最小化前處理的模型 AUC

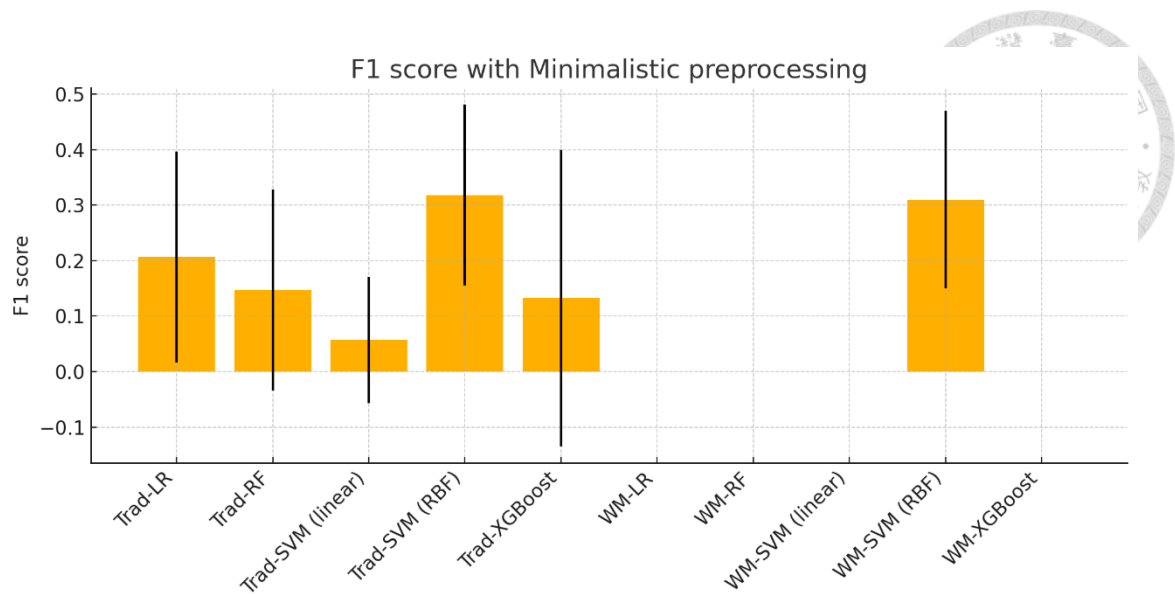
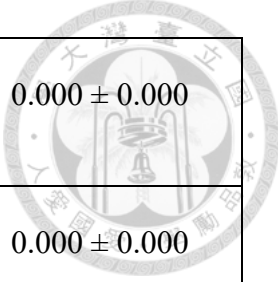


Figure 3.2-6: 使用最小化前處理的模型 F1 score

Table 3.2-4: 使用最小化前處理的模型結果一覽表

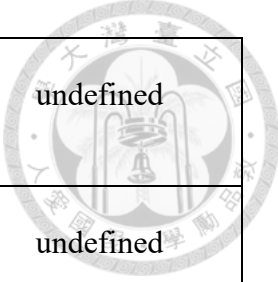
Minimalistic	Balanced accuracy	AUC	F1 score
Traditional Logistic regression	0.460 ± 0.145	0.487 ± 0.209	0.207 ± 0.190
Traditional Random forest	0.498 ± 0.067	0.548 ± 0.228	0.147 ± 0.181
Traditional SVM (linear)	0.360 ± 0.059	0.522 ± 0.144	0.057 ± 0.114
Traditional SVM (RBF)	0.517 ± 0.033	0.575 ± 0.303	0.318 ± 0.163
Traditional XGBoost	0.536 ± 0.111	0.595 ± 0.112	0.133 ± 0.267
WM Logistic regression	0.500 ± 0.000	0.519 ± 0.038	0.000 ± 0.000



WM Random forest	0.500 ± 0.000	0.500 ± 0.000	0.000 ± 0.000
WM SVM (linear)	0.500 ± 0.000	0.500 ± 0.000	0.000 ± 0.000
WM SVM (RBF)	0.500 ± 0.000	0.500 ± 0.000	0.310 ± 0.160
WM XGBoost	0.500 ± 0.000	0.500 ± 0.000	0.000 ± 0.000

Table 3.2-5: 使用最小化前處理的模型 balanced accuracy 及統計比較

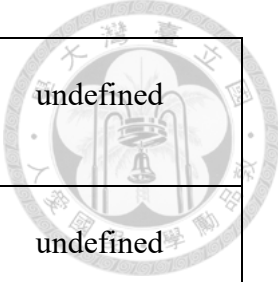
	Balanced accuracy	Shapiro-Wilk p-value	T-test (against 0.5) p-value
Traditional Logistic regression	0.460 ± 0.145	0.6562	0.6058
Traditional Random forest	0.498 ± 0.067	0.4278	0.9465
Traditional SVM (linear)	0.360 ± 0.059	0.4346	0.00864* (worse than 0.5)
Traditional SVM (RBF)	0.517 ± 0.033	0.00013*	0.3739 (t-test) 0.3173 (Wilcoxon)
Traditional XGBoost	0.536 ± 0.111	0.02699*	0.5538 (t-test) 0.6547 (Wilcoxon)
WM Logistic regression	0.500 ± 0.000 (all data being 0.5)	1	undefined



WM Random forest	0.500 ± 0.000 (all data are 0.5)	1	undefined
WM SVM (linear)	0.500 ± 0.000 (all data are 0.5)	1	undefined
WM SVM (RBF)	0.500 ± 0.000 (all data are 0.5)	1	undefined
WM XGBoost	0.500 ± 0.000 (all data are 0.5)	1	undefined

Table 3.2-6: 使用最小化前處理的模型 AUC 及統計比較

	AUC	Shapiro-Wilk p-value	T-test (against 0.5) p-value
Traditional Logistic regression	0.487 ± 0.209	0.2981	0.9090
Traditional Random forest	0.548 ± 0.228	0.7127	0.6973
Traditional SVM (linear)	0.522 ± 0.144	0.3844	0.7725
Traditional SVM (RBF)	0.575 ± 0.303	0.7205	0.6478
Traditional XGBoost	0.595 ± 0.112	0.5643	0.1633
WM Logistic regression	0.519 ± 0.038	0.00013	0.3739 (t-test) 0.3173 (Wilcoxon)



WM Random forest	0.500 ± 0.000 (all data are 0.5)	1	undefined
WM SVM (linear)	0.500 ± 0.000 (all data are 0.5)	1	undefined
WM SVM (RBF)	0.500 ± 0.000 (all data are 0.5)	1	undefined
WM XGBoost	0.500 ± 0.000 (all data are 0.5)	1	undefined

Chapter 4 Discussion

4.1 結論



由於過去臨床抗體腦炎預測量表皆無納入腦波發現，本研究主要目的在評估單次常規腦電圖，在疑似腦炎病患中，是否可作為獨立作為預測自體免疫腦炎（autoimmune encephalitis）抗體的參考。研究中比較了兩種不同的前處理策略、兩類特徵擷取方法、以及多種傳統機器學習分類演算法，但整體結果顯示，所有流程與搭配之模型的平衡正確率、AUC 等指標均無法顯著優於隨機猜測。以現有腦炎的腦波資料而言，搭配常見的腦波特徵擷取方式與傳統機器學習演算法，單一次臨床腦波的預測力相當有限，此說明腦波可能較適合做為輔助工具，而非獨立決策依據。這一點也側面驗證為何上述臨床量表並沒有腦波的成分。

方法學上，Automagic 流程原理上可以產出較無雜訊的腦波資料，搭配傳統特徵擷取與 linear SVM 時，預測抗體的平均能力稍稍優於隨機，但正確率仍然太低，尚且無法達到可資診斷參考的程度。最小化前處理對比之下成績更差，再次說明當前處理對訊號分析與機器學習的重要性。至於 WEASEL-MUSE，在(35)中，Guimarães 曾使用多個樣本數高的公開腦波資料集比較特徵擷取與機器學習方法，也發現 WEASEL-MUSE 表現不如預期，代表此方法可能不合適使用於腦波資料。

迄今為止，並沒有機器學習單獨使用腦波來預測抗體的存在，大多數的含腦波分析的論文都聚焦在預測抗體腦炎病患之預後(如(52))，或是需要搭配影像與臨床資料一起判斷(如(53))。總結而言，本研究之貢獻在於：確立單獨單次腦波無法獨立預測疑似腦炎病患之自體抗體、驗證去除雜訊在小樣本資料的重要性、比較不同訊號擷取方式對分析結果的影響、以及提供未來免疫腦炎腦波研究之比較基準。



4.2 研究限制

本研究的條件限制相當嚴苛，只有 46 位病患，其中只有 12 位確定抗體陽性，樣本數少且個數不平衡，處理上比較困難。因資料量少代表性亦可能不足，所以有待後續研究證實或反證抗體腦炎是否有特殊腦波變化。

另外，本研究只有採計第一次腦波，而腦炎病程可能起伏，最接近住院第一次的腦波特徵可能尚未顯露出來腦炎的樣貌，另外，癲癇用藥也可能會腦波變化，但本實驗並沒有收錄用藥狀況，且即便有收錄，樣本數小也很難進一步分析，此兩者皆為本實驗現階段無法避免的限制。

4.3 未來展望

未來相關研究可考慮的方向有：

1. 擴大收案樣本數：要驗證是否腦炎抗體病患之腦波與非抗體疑似腦炎之腦波是否有差異，需要更多的病患才有辦法達到。然而，因為疑似腦炎患者發生率並不高，以五年只有 46 位患者的收集速率，若要進一步發展需要多中心、前瞻性的合作計畫才較有機會。
2. 轉移學習(transfer learning)：以機器學習而言，因為直接增加抗體陰性一般癲癇病患只會加劇 class imbalance 所以並不合適。但是，若是以大量之癲癇病患來訓練 pre-trained model 再 transfer learning 可能有助於腦炎病患的資料分析。
3. 使用多模態資料：除卻腦波之外，加入臨床變項一起預測抗體存在性，不將



腦波視為單一資料來源。

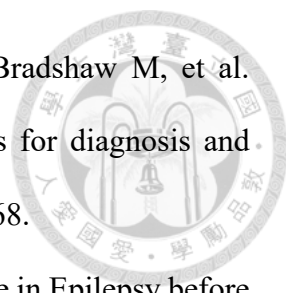
4. 做次分析，限縮探討範圍至 Anti-NMDAR 腦炎：一方面是目前僅有 Anti-NMDAR 腦炎有已知的腦波特色變化，另一方面是不同抗體腦炎之腦波可能特性不同，在個案數不多的前提下，與非抗體疑似腦炎的腦波特色可能觀察不到。另外，亦可以嘗試在這群病患嘗試 phase-amplitude coupling 看看是否 delta 徐波與上面疊加的快波之間有無關係。
5. 比較追蹤腦波：合併第一點，在擴大收案時確立腦波追蹤時程表，不論病患表現，都要不只做一次腦波，在更多的資料中探索是否非單次的腦波更能夠預測抗體存在。
6. 探索其他特徵擷取方式：由於腦波的特徵擷取方式相當多元，本實驗並沒有列如諸如 coherence 或是 phase-locking value (PLV) 等等特徵。
7. 探索腦波相關之 pre-trained neural network，以發掘腦波未被本研究發現，但可能還有的非線性隱藏特徵。

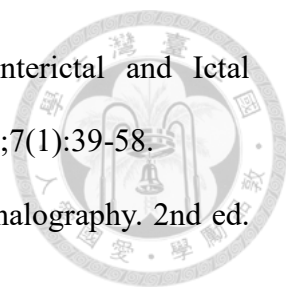
然而，下列第二至第六點實際上也綁在第一點上。本實驗只使用傳統特徵擷取方式與機器學習演算法，而不進行更多的特徵擷取方式或神經網路模型(neural network models, e.g. InceptionTime) 的最主要理由是個案數稀少。如果再加入擷取方式就會嚴重膨脹特徵維度，例如一個頻段的 pair-wise magnitude-squared coherence 就有 $21^2 = 441$ 個維度，加上之後本實驗的維度接近直接翻倍；而在如此小的樣本也不適合直接訓練神經網路模型，很容易造成過擬合(overfitting)。因此，若能擴大資料規模、針對特定抗體類型收案，並結合多模態資料，可能可望發展出更具實用價值的模型。

參考文獻



1. Graus F, Titulaer MJ, Balu R, Benseler S, Bien CG, Cellucci T, et al. A clinical approach to diagnosis of autoimmune encephalitis. *Lancet Neurol*. 2016;15(4):391-404.
2. Darnell RB, Posner JB. Paraneoplastic syndromes involving the nervous system. *New England Journal of Medicine*. 2003;349(16):1543-54.
3. Segal Y, Rotschild O, Mina Y, Maayan Eshed G, Levinson T, Paran Y, et al. Epidemiology of autoimmune encephalitis and comparison to infectious causes- Experience from a tertiary center. *Ann Clin Transl Neurol*. 2024;11(9):2337-49.
4. Shah S, Flanagan EP, Paul P, Smith CY, Bryant SC, Devine MF, et al. Population-Based Epidemiology Study of Paraneoplastic Neurologic Syndromes. *Neurol Neuroimmunol Neuroinflamm*. 2022;9(2).
5. Dalmau J, Graus F. Diagnostic criteria for autoimmune encephalitis: utility and pitfalls for antibody-negative disease. *Lancet Neurol*. 2023;22(6):529-40.
6. Lancaster E. Autoantibody Encephalitis: Presentation, Diagnosis, and Management. *J Clin Neurol*. 2022;18(4):373-90.
7. Vogrig A, Muniz-Castrillo S, Desestret V, Joubert B, Honnorat J. Pathophysiology of paraneoplastic and autoimmune encephalitis: genes, infections, and checkpoint inhibitors. *Ther Adv Neurol Disord*. 2020;13:1756286420932797.
8. Saraya AW, Worachotsueptrakun K, Vutipongsatorn K, Sonpee C, Hemachudha T. Differences and diversity of autoimmune encephalitis in 77 cases from a single tertiary care center. *BMC Neurol*. 2019;19(1):273.
9. Dubey D, Kothapalli N, McKeon A, Flanagan EP, Lennon VA, Klein CJ, et al. Predictors of neural-specific autoantibodies and immunotherapy response in patients with cognitive dysfunction. *Journal of neuroimmunology*. 2018;323:62-72.

- 
10. Abboud H, Probasco JC, Irani S, Ances B, Benavides DR, Bradshaw M, et al. Autoimmune encephalitis: proposed best practice recommendations for diagnosis and acute management. *J Neurol Neurosurg Psychiatry*. 2021;92(7):757-68.
11. Li Y, Tymchuk S, Barry J, Muppidi S, Le S. Antibody Prevalence in Epilepsy before Surgery (APES) in drug-resistant focal epilepsy. *Epilepsia*. 2021;62(3):720-8.
12. Dubey D, Singh J, Britton JW, Pittock SJ, Flanagan EP, Lennon VA, et al. Predictive models in the diagnosis and treatment of autoimmune epilepsy. *Epilepsia*. 2017;58(7):1181-9.
13. Chang YC, Nouri MN, Mirsattari S, Burneo JG, Budhram A. “Obvious” indications for Neural antibody testing in Epilepsy or Seizures: The ONES checklist. *Epilepsia*. 2022;63(7):1658-70.
14. Kandel ER, Koester J, Mack S, Siegelbaum S. Principles of neural science. Sixth edition. ed. New York: McGraw Hill; 2021. 1, 1646 pages p.
15. Greenfield LJ, Jr., Geyer JD, Carney PR. Reading EEGs : a practical approach. Place of publication not identified: LWW; 2012.
16. Blum AS, Rutkove SB. The clinical neurophysiology primer. Totowa, N.J: Humana Press; 2010. x, 526 pages : illustrations (some color) p.
17. Mostile G, Terranova R, Carlentini G, Contrafatto F, Terravecchia C, Donzuso G, et al. Differentiating neurodegenerative diseases based on EEG complexity. *Sci Rep*. 2024;14(1):24365.
18. Tomasello L, Carlucci L, Lagana A, Galletta S, Marinelli CV, Raffaele M, et al. Neuropsychological Evaluation and Quantitative EEG in Patients with Frontotemporal Dementia, Alzheimer's Disease, and Mild Cognitive Impairment. *Brain Sci*. 2023;13(6).
19. Davey Z, Gupta PB, Li DR, Nayak RU, Govindarajan P. Rapid Response EEG: Current State and Future Directions. *Curr Neurol Neurosci Rep*. 2022;22(12):839-46.

- 
20. Foldvary-Schaefer N, Grigg-Damberger MM. Identifying Interictal and Ictal Epileptic Activity in Polysomnograms. *Sleep Medicine Clinics*. 2012;7(1):39-58.
21. Daly DD, Pedley TA. *Current practice of clinical electroencephalography*. 2nd ed. New York: Raven Press; 1990. p. p.
22. Parvizi J, Spencer DC. *A Practical Approach to Stereo EEG*: Oxford University Press; 2021. 320 p.
23. Schmitt SE, Pargeon K, Frechette ES, Hirsch LJ, Dalmau J, Friedman D. Extreme delta brush: a unique EEG pattern in adults with anti-NMDA receptor encephalitis. *Neurology*. 2012;79(11):1094-100.
24. Parwani J, Ortiz JF, Alli A, Lalwani A, Ruxmohan S, Tamton H, et al. Understanding Seizures and Prognosis of the Extreme Delta Brush Pattern in Anti-N-Methyl-D-Aspartate (NMDA) Receptor Encephalitis: A Systematic Review. *Cureus*. 2021;13(9):e18154.
25. Simabukuro MM, Silva GDD, Castro LHM, Lucato LT. A critical review and update on autoimmune encephalitis: understanding the alphabet soup. *Arq Neuropsiquiatr*. 2022;80(5 Suppl 1):143-58.
26. Acharya JN, Hani A, Cheek J, Thirumala P, Tsuchida TN. American Clinical Neurophysiology Society Guideline 2: Guidelines for Standard Electrode Position Nomenclature. *J Clin Neurophysiol*. 2016;33(4):308-11.
27. Pedroni A, Bahreini A, Langer N. Automagic: Standardized preprocessing of big EEG data. *Neuroimage*. 2019;200:460-73.
28. Delorme A, Makeig S. EEGLAB: an open source toolbox for analysis of single-trial EEG dynamics including independent component analysis. *J Neurosci Methods*. 2004;134(1):9-21.
29. Bigdely-Shamlo N, Mullen T, Kothe C, Su KM, Robbins KA. The PREP pipeline:

standardized preprocessing for large-scale EEG analysis. *Front Neuroinform.* 2015;9:16.

30. Pion-Tonachini L, Kreutz-Delgado K, Makeig S. ICLabel: An automated electroencephalographic independent component classifier, dataset, and website. *Neuroimage.* 2019;198:181-97.

31. Salazar A, Vergara L. Independent component analysis (ICA) : algorithms, applications and ambiguities. New York: Nova Science Publisher's; 2018. viii, 334 pages p.

32. Harris CR, Millman KJ, van der Walt SJ, Gommers R, Virtanen P, Cournapeau D, et al. Array programming with NumPy. *Nature.* 2020;585(7825):357-62.

33. Virtanen P, Gommers R, Oliphant TE, Haberland M, Reddy T, Cournapeau D, et al. SciPy 1.0: fundamental algorithms for scientific computing in Python. *Nat Methods.* 2020;17(3):261-72.

34. Lee GR, Yamamoto R, Shou W, Kettimuthu R, Brigham E. PyWavelets: A Python package for wavelet analysis. *Journal of Open Source Software.* 2019;4(36):1237.

35. Guimarães LSA. Unraveling the Brain: A Quantitative Study of EEG Classification Techniques: Universidade de São Paulo; 2020.

36. Hjorth B. EEG analysis based on time domain properties. *Electroencephalogr Clin Neurophysiol.* 1970;29(3):306-10.

37. Merletti R, Lo Conte LR, De Luca CJ. A continuous wavelet based technique for the analysis of electromyography signals. *Medical & Biological Engineering & Computing.* 1999;37(5):512-9.

38. Schäfer P, Leser U, editors. Multivariate Time Series Classification with WEASEL+MUSE. Proceedings of the 23rd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD '17); 2017: ACM.

39. Schäfer P, editor SFA: A Symbolic Fourier Approximation and Index for Similarity

Search in High Dimensional Datasets. Proceedings of the 15th IEEE International Conference on Data Mining (ICDM); 2015: IEEE.

40. Abu-Mostafa YS, Magdon-Ismail M, Lin H-T. Learning from Data: A Short Course. New York: AMLBook; 2012.

41. Pedregosa F, Varoquaux G, Gramfort A, Michel V, Thirion B, Grisel O, et al. Scikit-learn: Machine Learning in Python. Journal of Machine Learning Research. 2011;12:2825–30.

42. Müller AC. Introduction to Machine learning with scikit-learn Cross Validation and Grid Search 2020 [Available from: <https://amueller.github.io/ml-training-intro/slides/03-cross-validation-grid-search.html>].

43. Kumar A. Python – Nested Cross Validation for Algorithm Selection 2020 [Available from: <https://vitalflux.com/python-nested-cross-validation-algorithm-selection/>].

44. Fawaz HI, Forestier G, Weber J, Idoumghar L, Muller P-A. InceptionTime: Finding AlexNet for Time Series Classification. Data Mining and Knowledge Discovery. 2020;34(6):1936–62.

45. Breiman L. Random Forests. Machine Learning. 2001;45(1):5–32.

46. Hastie T, Tibshirani R, Friedman JH. The elements of statistical learning : data mining, inference, and prediction. 2nd ed. New York, NY: Springer; 2009. xxii, 745 p. p.

47. Louppe G. Understanding Random Forests: From Theory to Practice [Ph.D. Thesis]. Liège, Belgium: University of Liège; 2014.

48. Chen T, Guestrin C, editors. XGBoost: A Scalable Tree Boosting System. Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD '16); 2016: ACM.

49. Ali ZH, Burhan AM. Hybrid machine learning approach for construction cost estimation: an evaluation of extreme gradient boosting model. Asian Journal of Civil

Engineering. 2023;24:2427–42.

50. Chawla NV, Bowyer KW, Hall LO, Kegelmeyer WP. SMOTE: Synthetic Minority Over-sampling Technique. *Journal of Artificial Intelligence Research*. 2002;16:321-57.

51. Obuchowski NA, Lieber ML, Wians FH. Sample Size Formulas For Estimating Areas Under the Receiver Operating Characteristic Curves With Precision and Assurance. *Statistics in Medicine*. 2004;23(9):1225-41.

52. Zhang Y, Shi X, Fan Z, Tu E, Wu D, Leng X, et al. Machine learning for the early prediction of long-term cognitive outcome in autoimmune encephalitis. *J Psychosom Res*. 2025;190:112051.

53. Musigmann M, Spiekers C, Stake J, Akkurt BH, Mora NGN, Sartoretti T, et al. Detection of antibodies in suspected autoimmune encephalitis diseases using machine learning. *Sci Rep*. 2025;15(1):10998.