

Tema 7. Combinación de modelos

- Algoritmo que usa varias instancias do mesmo **modelo base** (clasificador ou regresor): ensemble, comité
- Exemplo: combinación de árbores de clasificación/regresión
- O nome do algoritmo (meta-algoritmo), p.ex. bagging ou boosting, aplícase ao método usado para combinar os modelos
- Exemplo: grupo de clasificadores do mesmo tipo, entrenados de formas distintas, cunha votación para decidir a saída
- Os modelos base son usualmente **débiles**: non funcionan moi ben, pero son simples e de entrenamento rápido
- O obxectivo é que a súa combinación teña máis calidade que a suma dos clasificadores base individuais, dando un modelo **forte**
- Theodoridis: Machine learning: a Bayesian and optimization perspective, Academic Press, Cap. 7

Combinación de modelos (II)

- Os modelos base deben ser **diversos** (distintos entre eles) para que a súa combinación sexa forte (mellor que calquera dos modelos individuais)
- Aprenden distintas facetas do problema de clasificación
- A diversidade pódese deber a diferencias entre modelos base en:
 - 1) Inicialización do entrenamento: p.ex. combinación de redes neuronais MLP con distintas inicializacións aleatorias de pesos
 - 2) Sintonización dos hiper-parámetros: combinación de MLPs con distintos tamanos de capas ocultas
 - 3) Conxunto de entrenamento: distintos conxuntos de entrenamento para os distintos modelos base
- Os algoritmos de combinación pódense aplicar sobre distintos tipos de modelos base: bagging de árbores de decisión ou bagging de clasificadores KNN

Bagging (I)

- **Bootstrap aggregating:** varios modelos base entrénanse en varios conxuntos distintos do mesmo tamaño
- Creados usando o método de **bootstrap**: conxunto de patróns seleccionados aleatoriamente con repetición de patróns
- Na mostra seleccionada hai patróns seleccionados varias veces e outros non seleccionados nunca
- O uso de bootstrap diversifica os modelos base mostrándolles aos modelos distintos conxuntos de entrenamiento
- A saída decídese por votación (clasificación) ou promedio/ mediana (regresión) entre os modelos base
- Os modelos base son normalmente árbores de decisión, aínda que poden ser outros. En Weka podes cambiar o modelo (opción -W)

Bagging (II)

- Bootstrap mellora os modelos base porque reduce a súa varianza e o sobreaprendizaxe sen aumentar o sesgo
- As árbores tenden a sobreaprender o conxunto de entrenamiento
- Bagging fai que ás árbores sexan máis diversas porque entrenan sobre distintos conxuntos
- A votación/promedio non é tan sensíbel ao ruído nos datos e compensa esta sobreaprendizaxe das árbores individuais
- Hiper-parámetro B (bag size): nº de árbores. Sintonizábel, calidade pouco sensíbel a B se o nº de árbores é suficientemente elevado
- Función **bagging** en paquete **ipred** de R: B=25 por defecto. Normalmente, B~100-200, dependendo do tamaño dos datos

Bagging (III)

- Alternativamente a grid-search, pódese determinar B observando o “erro fora de bolso” (out-of-bag error, OOB):

Erro promedio, excluindo para cada árbore os patróns do seu conxunto de entrenamento (creado mediante bootstrap)

- O erro OOB estabilízase cando o nº B de árbores é dabondo
- Saída do bagging: se $z_b(\mathbf{x})$ é a saída do b -ésimo modelo base:

Clasificación:
$$z(\mathbf{x}) = \underset{c=1 \dots C}{\operatorname{arg\,max}} \left\{ \sum_{b=1}^B \delta[z_b(\mathbf{x}), c] \right\} \quad \delta(x, y) = \begin{cases} 1 & x = y \\ 0 & x \neq y \end{cases}$$

Regresión:
$$z(\mathbf{x}) = \frac{1}{B} \sum_{b=1}^B z_b(\mathbf{x})$$

Votación entre os
B clasificadores

Promedio das saídas
dos B regresores

Boosting

- Empuxa (boosts) cara arriba a calidade dos modelos base
- Como en bagging, os modelos base diferéncianse no conxunto de entrenamento: os seus patróns pondéranse de forma distinta para cada modelo base
- Modelo 1: tódolos patróns teñen pesos iguais
- Modelos posteriores: cada patrón pondérase en base ao erro cometido polos modelos anteriores nese patrón, comparado cos outros patróns (pesos w_i , con $i=1..N$)
- Pondéranse máis os patróns nos que houbo fallos para que sexan corrixidos polos seguintes modelos base
- A saída do comité $z(\mathbf{x})$ é un promedio ponderado das saídas dos modelos base, con pesos a_b , con $b=1..B$
- O máis popular: **adaboost** (adaptive boosting), clasificación

Adaboost (I)

$z_b(\mathbf{x}, \theta_b) \in \{\pm 1\}$: saída do clasificador b (coñecida)

- Clasificación binaria: $y_i, z(\mathbf{x}_i) \in \{\pm 1\}$: $z(\mathbf{x}) = \text{sign}[u(\mathbf{x})]$, $u = \sum_{b=1}^B a_b z_b(\mathbf{x}, \theta_b)$
- Combinación de B clasificadores $\{z_b(\mathbf{x}, \theta_b)\}_{b=1}^B$ con $z_b(\mathbf{x}, \theta_b) \in \{\pm 1\}$.

θ_b : parámetros entrenábeis do clasificador b

- Función de custe (a minimizar): $J[t, z(\mathbf{x})] = e^{-tz(\mathbf{x})}$

- Sexa $u_b(\mathbf{x})$ a suma parcial dos b primeiros clasificadores: $u_b(\mathbf{x}) = \sum_{k=1}^b a_k z_k(\mathbf{x}, \theta_k), b = 1, \dots, B$

- Os $u_b(\mathbf{x})$ calcúlanse iterativamente:

$$u_b(\mathbf{x}) = u_{b-1}(\mathbf{x}) + a_b z_b(\mathbf{x}, \theta_b), b = 2 \dots B$$

- Na iteración b (para o clasificador base b) calcúlanse a_b e θ_b :

$$(a_b, \theta_b) = \underset{a, \theta}{\operatorname{argmin}} \left\{ \sum_{i=1}^N w_i^b \exp[-y_i u_b(\mathbf{x}_i)] \right\}$$

- w_i^b é o peso de $\mathbf{x}_i$ na iteración b : $w_i^1 = 1/N$; $w_i^b = \exp[-y_i u_{b-1}(\mathbf{x}_i)]$, $b > 1$

Adaboost (II)

- Como $u_b(\mathbf{x}_i) = u_{b-1}(\mathbf{x}_i) + a_b z_b(\mathbf{x}_i, \theta_b)$:

$$(a_b, \theta_b) = \underset{a, \theta}{\operatorname{argmin}} \left\{ \sum_{i=1}^N w_i^b \exp \{ -y_i [u_{b-1}(\mathbf{x}_i) + a z_b(\mathbf{x}_i, \theta)] \} \right\}$$

- Iteración b: a mantiene constante e calcúlase θ_b (entrenamiento):

$$\theta_b = \underset{\theta}{\operatorname{argmin}} \left\{ \sum_{i=1}^N w_i^b \exp [-y_i a z_b(\mathbf{x}_i, \theta_b)] \right\}$$

- O entrenamiento persigue minimizar a suma de pesos dos patróns clasificados incorrectamente

$$\theta_b = \underset{\theta}{\operatorname{argmin}} \{ P_b(\theta) \} \quad P_b(\theta) = \sum_{y_i z_b(\mathbf{x}_i, \theta) < 0}^N w_i^b$$

$P_b(\theta)$ é a suma dos pesos dos patróns $\mathbf{x}_i$ nos que $y_i z_b(\mathbf{x}_i, \theta_b) < 0$: erros de clasificación do ensemble cos $b-1$ primeiros clasificadores base

- Coñecido θ_b , calcúlase a_b : $a_b = \underset{a}{\operatorname{argmin}} \{ P_b e^a + (1 - P_b) e^{-a} \}$

Adaboost (III)

- Derivando a expresión entre chaves e igualando a 0 obtense:

$$a_b = \frac{1}{2} \ln \left(\frac{1 - P_b}{P_b} \right)$$

- Coñecidos θ_b e a_b , os pesos w_i^{b+1} para $i=1, \dots, N$, calcúlanse como:

$$w_i^{b+1} = \frac{w_i^b \exp[-y_i a_b z_b(\mathbf{x}_i, \boldsymbol{\theta}_b)]}{Z_b}$$

onde Z_b e o factor de normalización: $Z_b = \sum_{i=1}^N w_i^b \exp[-y_i a_b z_b(\mathbf{x}_i, \boldsymbol{\theta}_b)]$

- O proceso continúa para o seguinte clasificador $b+1$ ata $b=B$
- Finalmente, coñecidos $\{a_b, \boldsymbol{\theta}_b\}_{b=1}^B$, a saída da combinación adaboost ven dada por:

$$z(\mathbf{x}) = \text{sign} \left[\sum_{b=1}^B a_b z_b(\mathbf{x}, \boldsymbol{\theta}_b) \right]$$

$$z_b(\mathbf{x}, \boldsymbol{\theta}_b) = \sum_{i=1}^{R_b} y_{bi} I_{bi}(\mathbf{x})$$

(Classification tree)

$$\boldsymbol{\theta}_b = \{R_b, \{y_{bi}, l_{bi}\}_{i=1}^{R_b}\}$$

Adaboost (IV)

Entrenamiento do clasificador base

$w_i^1 = 1/N, i=1..N$

for $b=1:B-1$ $P_b(\theta) = \sum_{y_i z_b(\mathbf{x}_i, \theta) < 0} w_i^b$; $\theta_b = \text{argmin}_{\theta} \{P_b(\theta)\}$; $a_b = \frac{1}{2} \log \left(\frac{1 - P_b^m}{P_b^m} \right)$

$P_b^m = P_b(\theta_b)$; $Z_b = 0$

for $i=1:N$

$w_i^{b+1} = w_i^b \exp[-y_i a_b z_b(\mathbf{x}_i, \theta_b)]$; $Z_b = Z_b + w_i^{b+1}$

endfor

for $i=1:N$

$w_i^{b+1} = w_i^{b+1} / Z_b$

endfor

endfor

Saída $z(\mathbf{x})$ para un patrón de teste $\mathbf{x}$: $z(\mathbf{x}) = \text{sign} \left[\sum_{b=1}^B a_b z_b(\mathbf{x}, \theta_b) \right]$

Random forest (I)

- Bosque aleatorio: combinación de árboles de clasificación ou regresión de modo que se corrixa a súa sobreaprendizaxe
- Usa bagging e selección aleatoria de características para o desdoblamento, deixando parte dos patróns fóra do entrenamiento
- Nunha árbore de decisión clásica, en cada nodo divídese a característica que máis reduce a entropía (clasificación) ou o erro cadrático medio (regresión)
- No bagging, se algunhas características son moi importantes, serán seleccionadas por case tódalas árbores: pouca diversidade
- No random forest, auméntase a diversidade usando, en cada árbore, un grupo de $q < n$ características escollidas aleatoriamente e distintas das seleccionadas polas outras árbores
- En cada nodo selecciónase a mellor característica neste grupo
- O entrenamiento é máis eficiente porque usa menos características

Random forest (II)

- Usualmente, o número q de entradas seleccionadas é $q=\sqrt{n}$ para clasificación e $q=n/2$ para regresión
- Se as saídas dos B modelos son variábeis aleatorias $\{z_b\}_{b=1}^B$, todos con varianza σ e correlación ρ , demóstrase que a varianza do RF é:

$$\text{var}\left(\frac{1}{B}\sum_{b=1}^B z_b\right)=\left(\frac{1-\rho}{B}+\rho\right)\sigma^2$$

- A selección aleatoria de características do RF:
 - 1) Incrementa lixeiramente o sesgo (inconvincente)
 - 2) Incrementa a varianza (σ^2) de cada árbore (inconvincente)
 - 3) Reduce a correlación (ρ) entre as árbores (vantaxe: +diversidade)
- A redución na correlación (3) é o termo máis importante: aumenta a calidade a respeito das árbores individuais

Random forest (III)

$RF = \emptyset$; $K = n^o$ máx. de nodos na árbore $\swarrow$ Árbore actual
for $b=1:B$ $\nwarrow$ N^o de nodos
 $S = \text{mostra de bootstrap de } \{\mathbf{x}_i, y_i\}_{i=1}^N$; $A = \emptyset$; $r=0$
 repeat
 $F = \{i_1 \dots i_q\} \subset \{1 \dots n\}$ con $q < n$, selección aleatoria
 Selecciona característica $j \in F$ e umbral $u \in V_j$ de modo que:
 clasificación: $(j, u) = \underset{i=1 \dots n, k \in V_i}{\arg \max} \{ \Delta E_{ik} \}$ $\swarrow$ $E = \text{entropía}$
 regresión: $(j, u) = \underset{i=1 \dots n; k \in V_i}{\arg \min} \left\{ \sum_{m \in A_1} (y_m - z_1)^2 + \sum_{m \in A_2} (y_m - z_2)^2 \right\}$
 $r = r + 1$
 Crea nodo n_r ($x_j < u$ e $x_j \geq u$) con j e u
 $A = A \cup \{n_r\}$
 $\swarrow$ $z_p(i, k) = \frac{1}{N_p} \sum_{q \in A_p} y_q, p=1,2$
 $A_1(i, k) = \{m : x_{mi} < k\}$
 $A_2(i, k) = \{m : x_{mi} \geq k\}$
 until $r > K$
 $RF = RF \cup A$
endfor
 Saída: clasificación: $z(\mathbf{x}) = \underset{c=1 \dots C}{\arg \max} \left\{ \sum_{b=1}^B \delta[z_b(\mathbf{x}), c] \right\}$; regresión: $z(\mathbf{x}) = \frac{1}{B} \sum_{b=1}^B z_b(\mathbf{x})$
 $\uparrow$ Votación $\delta(x, y) = \begin{cases} 1 & x = y \\ 0 & x \neq y \end{cases}$ $\uparrow$ Promedio

Random forest (IV)

- RF proporciona unha medida da importancia de cada característica
- Poucos hiperparámetros e pouco importantes (fácil de usar):
 - 1) N° de árbores B no bosque
 - 2) N° de características q a desdobrar en cada nodo
 - 3) N° mínimo de patróns para desdobrar un nodo interno
- Aumentando B non se incrementa o sobreaprendizaxe
- Moi paralelizábel
- Require pouco preprocesamento dos datos
- O uso de $q < n$ características ou de $N' < N$ patróns para o bootstrap é eficiente con datos grandes
- Moi bos resultados: método state-of-the-art