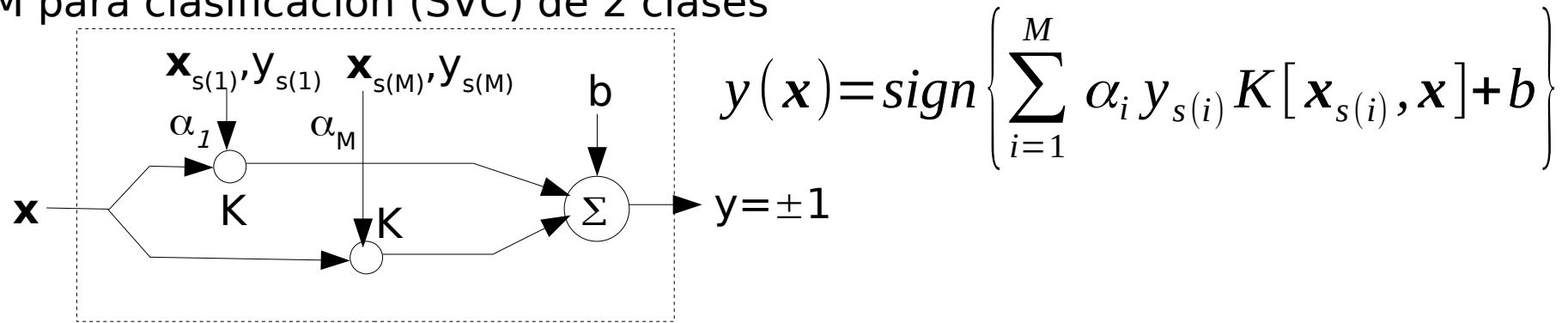


Tema 6. Máquinas de vectores de soporte (SVM)

SVM para clasificación (SVC) de 2 clases



- Clasificador lineal $y = \text{sgn}(\mathbf{w}^T \mathbf{x} + b)$ no espazo oculto. Saída $y = -1$ para unha clase, $y = +1$ para a outra; $\mathbf{x}_{s(1)} \dots \mathbf{x}_{s(M)}$ son M vectores de soporte
- Parámetros entrenábeis: $\{\alpha_i\}_{i=1}^M$, b ; hiper-parámetros sintonizábeis: regularización (λ), parámetro(s) do cerne (kernel) K
- A SVM combina dous ingredientes fundamentais:
 - 1) Uso dun cerne K para aumentar a separabilidade linear dos datos
 - 2) Selección do clasificador lineal que minimiza a sobreaprendizaxe

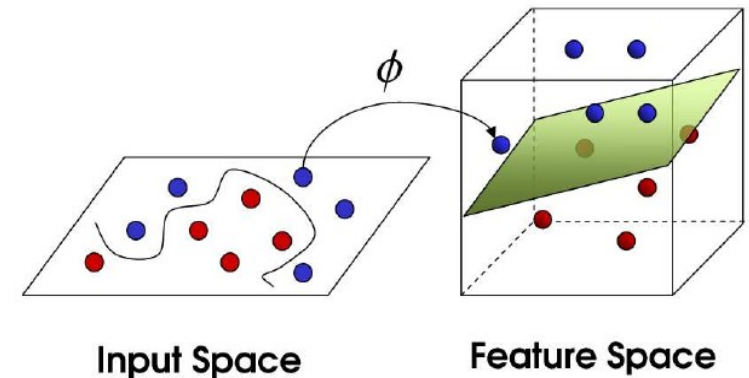
Cernes (I)

- Teorema de Cover (1965): a probabilidade de que N patróns no espazo n -dimensional sexan linearmente separábeis é 1 se $n \geq N-1$ e:

$$\frac{1}{2^N} \sum_{i=1}^n \binom{N-1}{i}$$

se $n < N-1$. Ésta é unha función crecente con n

- Aumentando a dimensionalidade n dos datos, aumenta a probabilidade de que sexan linearmente separábeis
- Aplicación cerne: mapea o espazo de entrada $\mathbb{R}^n$ a $\mathbb{R}^m$ con $m > n$:
 $\mathbf{x} \rightarrow \Phi(\mathbf{x})$
- A separabilidade linear é máis probábel neste segundo espazo, chamado espazo oculto ou de características
- Un cerne Φ usualmente cumpre que $\Phi(\mathbf{x})^\top \Phi(\mathbf{y}) = K(\mathbf{x}, \mathbf{y})$: produto escalar xeralizado. Se existe expresión para $\Phi(\mathbf{x})$ en $\mathbb{R}^m$, o cerne é separábel



Cernes (II)

- $K(\mathbf{x}, \mathbf{y})$ é unha medida de similaridade xeralizada entre $\mathbf{x}$ e $\mathbf{y}$
- A SVM é un clasificador linear de vector $\mathbf{w}$ no espazo oculto $\Phi(\mathbf{x})$
- Veremos que: $\mathbf{w} = \sum_{i=1}^M \alpha_i y_{s(i)} \Phi[\mathbf{x}_{s(i)}]$
sendo $M < N$ o nº de vectores de soporte $\{\mathbf{x}_{s(i)}\}_{i=1}^M$
- Polo tanto: $y = \text{sign}[\mathbf{w}^T \Phi(\mathbf{x}) + b] = \text{sign} \left\{ \sum_{i=1}^M \alpha_i y_{s(i)} \Phi[\mathbf{x}_{s(i)}]^T \Phi(\mathbf{x}) + b \right\}$
$$y(\mathbf{x}) = \text{sign} \left\{ \sum_{i=1}^M \alpha_i y_{s(i)} K[\mathbf{x}_{s(i)}, \mathbf{x}] + b \right\}$$
- Hai varios cernes distintos, con hiper-parámetros que se deben sintonizar para cada problema
- Cerne **gausiano** ou RBF: hai que sintonizar a anchura σ : Valores: $\{2^i\}_{-5}^{10}$
(espazo oculto de dimensión infinita)

$$K(\mathbf{x}, \mathbf{y}) = \exp \left(\frac{-|\mathbf{x} - \mathbf{y}|^2}{2\sigma^2} \right)$$

Cernes (III)

- Cerne polinómico $K(\mathbf{x}, \mathbf{y}) = (\mathbf{x}^T \mathbf{y} + a)^b$. Hiperparámetros:
 - 1) a : sesgo, $-\mu \leq a \leq \mu$ se $\mathbf{x}, \mathbf{y}$ están estandarizados: $|\mathbf{x}^T \mathbf{y}| < \mu$
 - 2) b : grao=1,2,3
- Hai outros cernes, p.ex. para cadeas de caracteres
- Sen cerne, usas un cerne linear $K(\mathbf{x}, \mathbf{y}) = \mathbf{x}^T \mathbf{y}$, clasificador linear no espazo de entrada, sen hiper-parámetros
- O cerne que mellor funciona normalmente é o gausiano, que so ten un hiperparámetro: a anchura σ
- A calidade da SVM (acerto en clasificación, RMSE en regresión) mostra un extremo para o mellor valor de σ
- A SVM con cerne linear funciona bastante peor que a SVM gausiana, agás con n moi alto, porque a separabilidade linear xa é moi elevada no espazo de entrada

Minimización da sobreaprendizaxe

- Teoría da aprendizaxe estatística (V. Vapnik, anos 60)
- Dimensión de Vapnik-Chervonenkis (h) dun clasificador de 2 clases é o máximo número de patróns que pode aprender sen erros, sexa cal sexa o etiquetado de clases dos patróns
- Mide a complexidade do clasificador: canto mais alta sexa h , máis sobreaprendizaxe. Débese minimizar
- Para un **clasificador linear** nun espazo n -dimensional $\mathbb{R}^n$: $h=n+1$
- Se os patróns verifican $|\mathbf{x}| < D$ e ρ é o **marxe** (mínima distancia entre $\mathbf{x}$ e o hiperplano que separa entre clases), entón:

$$h \leq \min \left(\left\lceil \frac{D}{\rho^2} \right\rceil, n \right) + 1 \quad \text{Para un clasificador linear}$$

- Débese maximizar a marxe ρ para minimizar a cota superior de h , e polo tanto a sobreaprendizaxe (risco estrutural)

Saída da SVM (I)

Distancia entre $\Phi(\mathbf{x}_i)$ e o plano

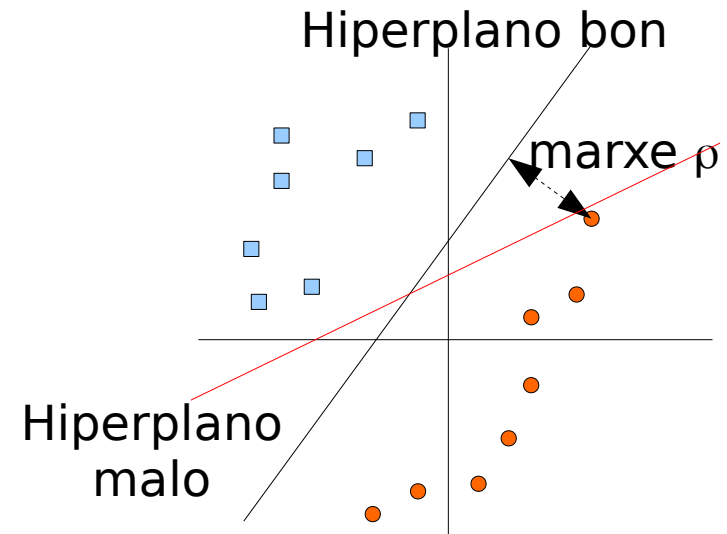
- A marxe ρ no espazo oculto é: $\rho = \min_{i=1 \dots N} \frac{|\mathbf{w}^T \Phi(\mathbf{x}_i) + b|}{|\mathbf{w}|}$

- Impónse a condición $|\mathbf{w}^T \Phi(\mathbf{x}_i) + b| = 1$, para o $\mathbf{x}_i$ máis cercano, así que $\rho = 1/|\mathbf{w}|$

- O erro de entrenamento avalíase como $\xi_i = \max(0, 1 - y_i(\mathbf{w}^T \Phi(\mathbf{x}_i) + b))$ para $\Phi(\mathbf{x}_i)$

- $\xi_i > 0 \leftrightarrow \Phi(\mathbf{x}_i)$ clasifícase mal ($\xi_i > 1$) ou ben pero moi cerca de estar mal ($0 < \xi_i < 1$)

- Hai que maximizar (λ =hiperparámetro de regularización):



Erro de entrenamento (risco empírico)

Erro de xeralización (risco estrutural)

$$J(\mathbf{w}, b, \vec{\xi}) = |\mathbf{w}|^2 + \lambda \sum_{i=1}^N \xi_i$$

coas condicións: $\xi_i \geq 0, y_i(\mathbf{w}^T \Phi(\mathbf{x}_i) + b) \geq 1 - \xi_i, i = 1 \dots N$

Saída da SVM (II)

- Método dos multiplicadores de Lagrange: $\{\alpha_i, \beta_i\}_{i=1}^N$

$$L(\mathbf{w}, b, \vec{\xi}, \vec{\beta}, \vec{\alpha}) = |\mathbf{w}|^2 + \lambda \sum_{i=1}^N \xi_i + \sum_{i=1}^N \beta_i \xi_i + \sum_{i=1}^N \alpha_i [y_i (\mathbf{w}^T \vec{\Phi}(\mathbf{x}_i) + b) - 1 + \xi_i]$$

- Derivando a respecto de $\mathbf{w}$ e igualando a $\mathbf{0}$: $\mathbf{w} = \sum_{i=1}^N \alpha_i y_i \vec{\Phi}(\mathbf{x}_i)$
- Vector $\mathbf{w}$: combinación linear dos patróns de entrenamento. Isto non é escalábel para N grandes
- Solución: como veremos, $\alpha_i = 0$ para moitos i (solución dispersa)
- Derivando a respecto de b, ξ_i , α_i e β_i transfórmase o problema a

maximizar $\vec{\alpha}^T \mathbf{1} - \frac{\vec{\alpha} \mathbf{K} \vec{\alpha}}{2}$, sendo $K_{ij} = K(\mathbf{x}_i^T, \mathbf{x}_j)$, coas condicións:

$$\vec{\alpha}^T \mathbf{y} = 0, 0 \leq \alpha_i \leq \lambda, \beta_i \xi_i = 0, \alpha_i \left\{ y_i \left[\sum_{j=1}^N \alpha_j K(\mathbf{x}_j^T, \mathbf{x}) + b \right] \right\} = 0, i = 1, \dots, N$$

Saída da SVM (III)

- Este problema de optimización resólvese usando cálculo numérico iterativo.
- Canto máis difícil é o problema, máis tempo tarda en resolverse a optimización. Pode ser máis rápida en problemas máis grandes
- Só se requiren M patróns de entrenamiento (vectores de soporte) $\mathbf{x}_{s(1)} \dots \mathbf{x}_{s(M)}$, para os que $0 \leq \alpha_i \leq \lambda$, sendo $\alpha_i = 0$ para os restantes patróns (solución dispersa):

$$\mathbf{w} \text{ no espazo oculto} \longrightarrow \mathbf{w} = \sum_{i=1}^M \alpha_i y_{s(i)} \Phi[\mathbf{x}_{s(i)}]$$

Cálculo do sesgo

- Sendo $(\mathbf{x}_m, y_m)$ un vector de soporte: $b = y_m - \sum_{i=1}^M \alpha_i y_{s(i)} K[\mathbf{x}_m, \mathbf{x}_{s(i)}]$
- Sustituíndo $\mathbf{w}$ en $z = \text{sign}(\mathbf{w}^T \Phi(\mathbf{x}) + b)$ e usando que $\Phi(\mathbf{x})^T \Phi(\mathbf{y}) = K(\mathbf{x}, \mathbf{y})$ obtemos a expresión da SVM:

$$y(\mathbf{x}) = \text{sign} \left\{ \sum_{i=1}^M \alpha_i y_{s(i)} K[\mathbf{x}_{s(i)}, \mathbf{x}] + b \right\}$$

Hiperparámetros e clasificación multi-clase (I)

- **Hiperparámetros sintonizáveis**, cerne gaussiano, grid-search:

λ (regularización): $2^{-5}..2^{15}$: os resultados non son moi sensíbeis ao seu valor, é pouco importante

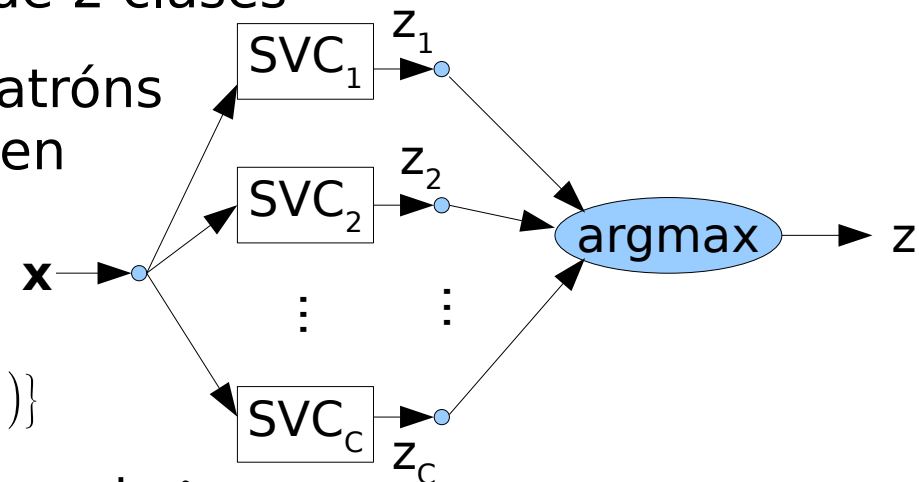
σ (ancho do cerne gaussiano): $2^{-3}..2^7$: moi importante nos resultados: o mellor valor está a miúdo no medio dos extremos

- **Clasificación multiclase**: se $C > 2$ é elevado, usar o esquema one-vs-all (**OVA**), de $O(C)$, onde a SVM i -ésima clasifica os patróns entre a clase i e as restantes, usa C SVMs de 2 clases

- A SVM i -ésima entrena con tódolos patróns ($y_k = 1$ se $\mathbf{x}_k$ pertence á clase i , $y_k = -1$ en caso contrario)

$$z_i(\mathbf{x}) = \sum_{j=1}^M \alpha_j y_{s(j)} K[\mathbf{x}_{s(j)}, \mathbf{x}] + b$$

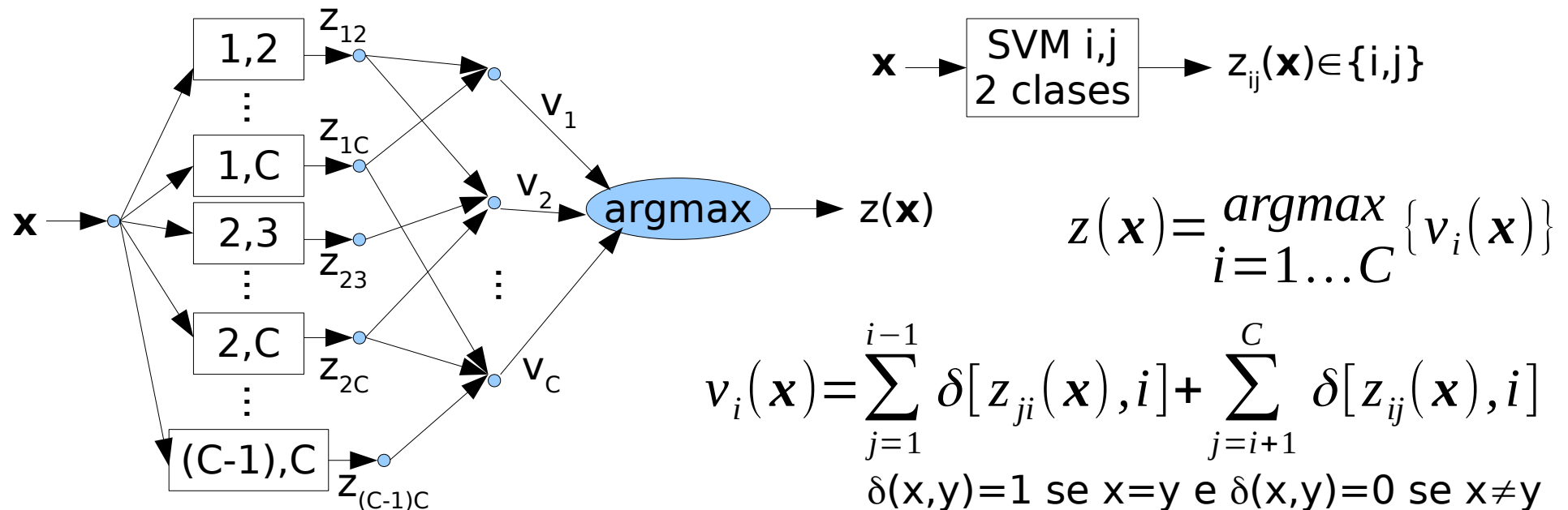
$$z(\mathbf{x}) = \underset{i=1 \dots C}{\operatorname{argmax}} \{z_i(\mathbf{x})\}$$



- Todas as SVMs usan os mesmos valores de λ e σ

Clasificación multi-clase (II)

- Con $C > 2$ baixo, emprégase un esquema one-vs-one (**OVO**). Usa $C(C-1)/2$ SVMs de 2 clases. Acerta máis que OVA. É máis lento que OVA porque ten complexidade $O(C^2)$
- A SVM ij clasifica entre a clase i e a j , con $i=1..C-1$ e $j=i+1..C$, entrenando so cos patróns $\mathbf{x}_k$ destas clases ($y_k=1$ se $\mathbf{x}_k$ é da clase i e $y_k=-1$ se $\mathbf{x}_k$ é da clase j)



Regresión regularizada con cernes (kernel ridge regression)

- Sexa un cerne $K: \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$, p. ex. gaussiano: $K(\mathbf{x}, \mathbf{y}) = \exp\left(\frac{-|\mathbf{x} - \mathbf{y}|^2}{2\sigma^2}\right)$
- Regresión non linear: modela a saída como:

$$f(\mathbf{x}) = \sum_{i=1}^N w_i K(\mathbf{x}, \mathbf{x}_i) = \mathbf{w}^T \mathbf{k}(\mathbf{x}) \quad \mathbf{k}(\mathbf{x}) = [K(\mathbf{x}, \mathbf{x}_1) \dots K(\mathbf{x}, \mathbf{x}_N)]^T$$

- A función K é un cerne (producto escalar xeralizado):

$$\mathbf{w} = \underset{\mathbf{v}}{\operatorname{argmin}} \{J(\mathbf{v})\} \quad J(\mathbf{w}) = \sum_{i=1}^N \left[y_i - \sum_{j=1}^N w_j K(\mathbf{x}_i, \mathbf{x}_j) \right]^2 + \lambda \|\mathbf{w}\|^2$$

- $J(\mathbf{w})$ pode describirse como $J(\mathbf{w}) = (\mathbf{y} - \mathbf{K}\mathbf{w})^T (\mathbf{y} - \mathbf{K}\mathbf{w}) + \lambda \mathbf{w}^T \mathbf{K}^T \mathbf{w}$, onde $\mathbf{K}$ é a matriz $N \times N$ de cernes con $K_{ij} = K(\mathbf{x}_i, \mathbf{x}_j)$. Derivando e igualando a $\mathbf{0}$:

$$(\mathbf{K}^T \mathbf{K} + \lambda \mathbf{I}) \mathbf{w} = \mathbf{K}^T \mathbf{y} \rightarrow \mathbf{w} = (\mathbf{K} + \lambda \mathbf{I})^{-1} \mathbf{y} \quad \text{— } \mathbf{I} \text{ é a matriz identidade}$$

- Saída para un patrón de teste $\mathbf{x}$: $z = \mathbf{y}^T (\mathbf{K} + \lambda \mathbf{I})^{-1} \mathbf{k}(\mathbf{x})$

Support vector regression ε -SVR (I)

- O RMSE non é moi axeitado para regresión porque os artefactos (valores moi alonxados da media) dan diferencias cadráticas altas

- Úsase a función linear ε -insensitiva: $\xi_i = \max(0, |y_i - z(\mathbf{x}_i)| - \varepsilon)$

- Minimizar a función de custe: $J = \frac{|\mathbf{w}|^2}{2} + \lambda \sum_{i=1}^N (\xi_i + \xi_i')$

coas condicións:

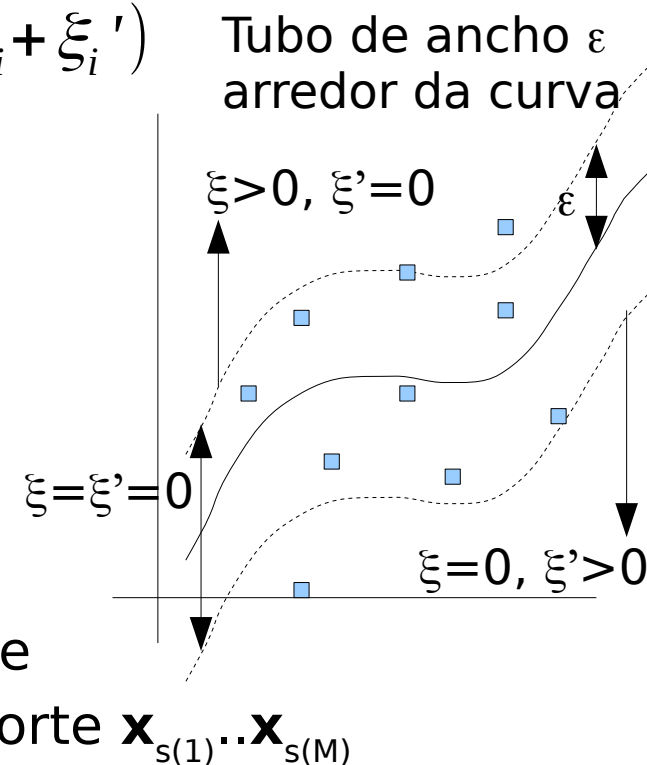
$$y_i - \mathbf{w}^T \Phi(\mathbf{x}_i) - b \leq \varepsilon + \xi_i', i = 1 \dots N$$

$$\mathbf{w}^T \Phi(\mathbf{x}_i) + b - y_i \leq \varepsilon + \xi_i, i = 1 \dots N$$

$$\xi_i \geq 0, \xi_i' \geq 0, i = 1 \dots N$$

- Solución: $\mathbf{w} = \sum_{i=1}^M (\alpha_i - \alpha_i') \Phi[\mathbf{x}_{s(i)}]$

sendo α_i e α_i' os M multiplicadores de Lagrange asociados ás condicións e aos vectores de soporte $\mathbf{x}_{s(1)} \dots \mathbf{x}_{s(M)}$



Support vector regression ε -SVR (II)

- Cálculo do sesgo: $b = \varepsilon - y_m + \sum_{i=1}^M (\alpha_i - \alpha_i') K[\mathbf{x}_m, \mathbf{x}_{s(i)}]$

onde $\mathbf{x}_m$ é un vector de soporte e y_m é a súa saída desexada

- Saída da ε -SVR: $z(\mathbf{x}) = b + \sum_{i=1}^M (\alpha_i - \alpha_i') K[\mathbf{x}, \mathbf{x}_{s(i)}]$
- Hiperparámetros sintonizábeis con cerne gausiano: igual que en clasificación: regularización (λ) e ancho do cerne (σ)
- Os vectores de soporte $\mathbf{x}_i$ verifican $z(\mathbf{x}_i) \in [y_i - \varepsilon, y_i + \varepsilon]$, e $0 \leq \alpha_i, \alpha_i' \leq \lambda$. Os restantes $\mathbf{x}_i$ verifican $\alpha_i = \alpha_i' = 0$. O valor de ε non é sensíbel ($\varepsilon = 0.1$ funciona ben)
- A ε -SVR é o mellor algoritmo de regresión que existe

Complexidade e implementacións

- O entrenamento da SVM é unha optimización cadrática con complexidade $O(N^3)$ e memoria $O(N^2)$; OVA é $O(C^2)$, OVO é $O(C)$
- Implementacións eficientes: $O(N^p)$ con $1 \leq p \leq 2.3$
- Lenta para $N > 10.000-50.000$, dependendo do nº de entradas
- Con patróns moi anchos (n elevado), usa cerne linear porque non é necesario mapear a un espazo multi-dimensional

$$y(\mathbf{x}) = \text{sign}(\mathbf{w}^T \mathbf{x} + b), \mathbf{w} = \sum_{i=1}^M \alpha_i y_{s(i)} \mathbf{x}_{s(i)}, b = y_m - \sum_{i=1}^M \alpha_i y_{s(i)} \mathbf{x}_m^T \mathbf{x}_{s(i)} \quad \text{Clasificación}$$

$(\mathbf{x}_m, y_m)$: vector de soporte

$$y(\mathbf{x}) = \mathbf{w}^T \mathbf{x} + b, \mathbf{w} = \sum_{i=1}^M (\alpha_i - \alpha_i') \mathbf{x}_{s(i)}, b = \varepsilon - y_m + \sum_{i=1}^M (\alpha_i - \alpha_i') \mathbf{x}_m^T \mathbf{x}_{s(i)} \quad \text{Regresión}$$

- Os cernes linear e gausiano dan calidades similares con n alto
- **LibSVM**: accesíbel dende C++, Octave/Matlab, Python, Weka/Java. Proporciona SVC multi-clase OVO, ε -SVR, e outras
- Función **ksvm** en paquete **kernlab** de R