

Introducción a las Máquinas de Soporte Vectorial

Estimación de funciones

El problema general que intentaremos abordar aquí es el siguiente:

Mediante algún proceso observacional/experimental se colectaron un conjunto de pares de datos (nuestra base de hechos) $F = \{(\mathbf{x}_i, y_i)\}_{i=1}^N$ donde $\mathbf{x}_i = \{x_i^1, \dots, x_i^p\} \in \mathcal{R}^p$ es un vector **aleatorio** de entradas del sistema bajo estudio y y_i la salida, univariada, correspondiente a dicha entrada. Denominaremos a este conjunto F como “conjunto de entrenamiento” o nuestra base de hechos (Base de datos).

El objetivo es intentar entender y modelar el sistema que, ante una entrada $\mathbf{x}_i$ produce una salida y_i , esto quiere decir que existe una f tal que $y_i = f(\mathbf{x}_i)$, pero esta función f es desconocida. Esto implica encontrar una función $\hat{f}(\mathbf{x}, \boldsymbol{\beta})$, con $\boldsymbol{\beta} \in \mathcal{B}$ (espacio de parámetros) de manera tal que minimizamos el siguiente funcional

$$R(\boldsymbol{\beta}) = \int L(y, \hat{f}(\mathbf{x}, \boldsymbol{\beta})) dF(\mathbf{x}, y) \quad (1.1)$$

Donde $L()$ se denomina función de pérdida y dependiendo del tipo de variable de respuesta (y_i) podemos definir al menos dos tipos básicos de problemas a resolver:

1. Reconocimiento de patrones (Clasificación): En este caso la salida del sistema desconocido puede tomar solo 2 valores $y_i = \{0, 1\}$ y sea $\hat{f}(\mathbf{x}, \boldsymbol{\beta})$, con $\boldsymbol{\beta} \in \mathcal{B}$ un conjunto de funciones “indicadoras” que solo puede tomar los valores 0 y 1. Consideremos la siguiente función de pérdida:

$$L(y, \hat{f}(\mathbf{x}, \boldsymbol{\beta})) = \begin{cases} 0 & \text{si } y = \hat{f}(\mathbf{x}, \boldsymbol{\beta}) \\ 1 & \text{si } y \neq \hat{f}(\mathbf{x}, \boldsymbol{\beta}) \end{cases} \quad (1.2)$$

Para esta función de pérdida, el función 1.1 determina la probabilidad de error de clasificación.

2. Regresión (Predicción, aproximación de funciones)

Sea $y_i \in \mathcal{R}$ y sea $\hat{f}(\mathbf{x}, \boldsymbol{\beta})$ un conjunto de funciones reales que contiene la siguiente función de regresión $\hat{f}(\mathbf{x}, \boldsymbol{\beta}_0) = \int y dF(y|\mathbf{x})$

Se sabe que dicha función de regresión es la que minimiza el funcional 1.1 con la siguiente

$$\text{función de pérdida } L(y, \hat{f}(\mathbf{x}, \boldsymbol{\beta})) = (y - \hat{f}(\mathbf{x}, \boldsymbol{\beta}))^2 \quad (1.3)$$

En términos generales podemos describir el problema de la siguiente manera: Sea la medida de probabilidad $F(\mathbf{z})$ definida en el espacio $\mathcal{Z}$. Considere el conjunto de funciones $Q(\mathbf{z}, \boldsymbol{\beta})$ con $\boldsymbol{\beta} \in \mathcal{B}$. El objetivo es minimizar el funcional de riesgo $R(\boldsymbol{\beta}) = \int Q(\mathbf{z}, \boldsymbol{\beta}) dF(\mathbf{z})$ con $\boldsymbol{\beta} \in \mathcal{B}$. donde $F(\mathbf{z})$ es desconocida pero contamos con un conjunto de muestras i.i.d $\mathbf{z}_1, \dots, \mathbf{z}_n$.

Minimización del riesgo empírico (MRE)

Con el objeto de minimizar el funcional 1.3 bajo una función de distribución $F(\mathbf{z})$ desconocida, podemos aplicar el siguiente principio inductivo:

- i. El funcional de riesgo $R(\boldsymbol{\beta})$ se reemplaza por el funcional de riesgo empírico $R_{emp}(\boldsymbol{\beta}) = \frac{1}{N} \sum_{i=1}^N Q(\mathbf{z}_i, \boldsymbol{\beta})$ (1.4) construido en base al conjunto de entrenamiento F .
- ii. Se aproxima la función $Q(\mathbf{z}, \boldsymbol{\beta}_0)$ que minimiza el riesgo 1.1 mediante la función $Q(\mathbf{z}, \boldsymbol{\beta})$ que minimiza el riesgo empírico.

Este concepto, MRE, es bastante general pudiendo considerarse a los métodos de mínimos cuadrados (MC) y de máxima verosimilitud (MV) realizaciones del MRE bajo funciones de pérdida específicas (la 1.2) en el primer caso en $L(p(\mathbf{x}, \boldsymbol{\beta})) = -\log(p(\mathbf{x}, \boldsymbol{\beta}))$. Con “p” función de densidad de probabilidad (conocida en el caso de MV)

Operar mediante el principio de MRE tiene la ventaja de que:

1. No depende de $F(\mathbf{z})$
2. En teoría es posible minimizar 1.4 con respecto a $\boldsymbol{\beta}$

Sea $\hat{\boldsymbol{\beta}}$ el vector de parámetros que minimiza $R_{emp}(\boldsymbol{\beta})$ a través de $\hat{f}(\mathbf{x}, \boldsymbol{\beta})$ y sea $\boldsymbol{\beta}_0$ el vector de parámetros que minimiza a $R(\boldsymbol{\beta})$ a través de $\hat{f}(\mathbf{x}, \boldsymbol{\beta}_0)$ con $\boldsymbol{\beta}_0$ y $\hat{\boldsymbol{\beta}} \in \mathcal{B}$ entonces:

¿Bajo qué condiciones podemos decir que $\hat{f}(\mathbf{x}, \hat{\boldsymbol{\beta}})$ esta “cerca” de la función deseada $\hat{f}(\mathbf{x}, \boldsymbol{\beta}_0)$ medida a través de la discrepancia entre $R_{emp}(\boldsymbol{\beta})$ y $R(\boldsymbol{\beta})$?

Dado $\boldsymbol{\beta} = \boldsymbol{\beta}^*$ “fijo”, el riesgo funcional $R(\boldsymbol{\beta}^*)$ determina la “Esperanza Matemática” de una variable aleatoria $Q(\boldsymbol{\beta}^*) = L(y, \hat{f}(\mathbf{x}, \boldsymbol{\beta}))$, mientras que el riesgo empírico $R_{emp}(\boldsymbol{\beta}^*)$ es la media muestral de $Q(\boldsymbol{\beta}^*)$, y por la ley de los grandes números tenemos que

$$R_{emp}(\boldsymbol{\beta}^*) \xrightarrow{N \rightarrow \infty} R(\boldsymbol{\beta}^*) \quad (1.5)$$

Lo que constituye la razón teórica para el uso del funcional empírico 1.4.

Sin embargo, la sola convergencia del valor esperado no implica que el mínimo de $R_{emp}(\boldsymbol{\beta})$ sea el mínimo de $R(\boldsymbol{\beta})$.

Para que esto suceda debemos imponer ciertas restricciones:

1. Que el $R_{emp}(\boldsymbol{\beta})$ aproxime uniformemente a $R(\boldsymbol{\beta})$ con una precisión ε tal que $|R_{emp}(\hat{\boldsymbol{\beta}}) - R(\boldsymbol{\beta}_0)| < 2\varepsilon \Rightarrow \forall \boldsymbol{\beta} \in \mathcal{B} \text{ y } \varepsilon > 0 \text{ se tiene que}$
$$P(\sup_{\boldsymbol{\beta}} |R(\boldsymbol{\beta}) - R_{emp}(\boldsymbol{\beta})| > \varepsilon) \xrightarrow{N \rightarrow \infty} 0$$

Entonces, podemos ahora definir formalmente el principio Minimización del Riesgo Empírico [Vapnik]:

1. En lugar del funcional de riesgo $R(\boldsymbol{\beta})$, construimos el funcional de riesgo empírico como

$R_{emp}(\beta) = \frac{1}{N} \sum_{i=1}^N L(y_i, \hat{f}(x_i, \beta))$ sobre la base de muestras i.i.d $\{x_i, y_i\}$, con $i=1,2,\dots,N$

2. Sea $\hat{\beta}$ el vector de parámetros que minimiza el funcional de riesgo empírico $R_{emp}(\hat{\beta})$ sobre el espacio de parámetros $\mathcal{B}$. Luego el $R_{emp}(\hat{\beta})$ converge en probabilidad al valor mínimo posible del funcional $R(\beta)$ verdadero, con $\beta \in \mathcal{B}$, a medida que el tamaño N de muestras del conjunto S se hace infinitamente grande, dado que el funcional de riesgo empírico $R_{emp}(\hat{\beta})$ converge uniformemente a $R(\beta)$.
3. Definimos convergencia uniforme mediante $P(\sup_{\beta} |R(\beta) - R_{emp}(\beta)| > \varepsilon) \xrightarrow{N \rightarrow \infty} 0$
Como una condición necesaria y suficiente para la consistencia del principio de minimización del riesgo empírico.

El principio de MRE en clasificación

En este caso tenemos que la salida “y” solo puede tomar dos valores posibles $\{0,1\}$. Por lo tanto la función de pérdida será

$$L(y, \hat{f}(x, \beta)) = \begin{cases} 0 & \text{si } y = \hat{f}(x, \beta) \\ 1 & \text{si } y \neq \hat{f}(x, \beta) \end{cases} \quad (1.6)$$

es decir que el valor 1 de L implica un “error”, por ende

$$R_{emp}(\beta) = \frac{1}{N} \sum_{i=1}^N L(y_i, \hat{f}(x_i, \beta))$$

es el error de entrenamiento o frecuencia de ocurrencia de un error (error cometido sobre el conjunto de entradas utilizado para encontrar el parámetro $\hat{\beta}$) y el funcional

$$R(\beta) = \int L(y, \hat{f}(x, \beta)) dF(x, y)$$

es la probabilidad de error de clasificación, denotada por $P(\beta)$. Luego por la ley de los grandes números tenemos que la frecuencia empírica de ocurrencia de un evento (error) converge casi seguramente a la $P(\text{error})$ a medida que $N \rightarrow \infty$. Como vimos más arriba, tenemos que con probabilidad α se cumple que $\sup_{\beta} |R(\beta) - R_{emp}(\beta)| \geq \varepsilon \Rightarrow$ con $p=1-\alpha$

$$P(\beta) = R(\beta) < R_{emp}(\beta) + \varepsilon(N, VC, \alpha, R_{emp}(\beta)) \quad (1.7)$$

Esto nos quiere decir que, si utilizamos un algoritmo que encuentre el $\hat{\beta}$ utilizando el principio MRE, el error de dicha función está acotado por el error empírico y un término que depende de la cantidad de muestras N, de la precisión α y de un parámetro “VC” conocido como dimensión de Vapnik-Chervonenkis (dimensión VC).

La dimensión VC juega un rol fundamental en la teoría del aprendizaje estadístico de donde se deriva el algoritmo de Máquinas de Soporte Vectorial. La dimensión VC es un concepto puramente combinatorial y no geométrico. La dimensión VC es un número que indica la cantidad de veces que un conjunto de datos F puede partirse en dos conjuntos disjuntos ($\mathcal{L}_0$ y $\mathcal{L}_1$) de manera tal que para cada una de estas particiones se satisfaga la siguiente ecuación:

$$f(x, \beta) = \begin{cases} 0 & \text{si } x \in \mathcal{L}_0 \\ 1 & \text{si } x \in \mathcal{L}_1 \end{cases} \quad (1.8)$$

sin errores.

La ecuación 1.7 junto con la definición de la dimensión VC nos provee de varios conceptos teóricos que nos ayudan a encontrar una mejor solución en pos de minimizar la $P(error)$.

En busca de la solución, el hiperplano óptimo para patrones separables linealmente.

Supongamos que contamos con un conjunto de entrenamiento $F = \{(\mathbf{x}_i, y_i)\}_{i=1}^N$ donde $\mathbf{x}_i$ es la i -ésima muestra con $y_i \in \{-1, 1\}$ como su salida deseada (que sea 0 o -1 es indistinto, pero continuaremos con -1 por simpleza del análisis que sigue). Por el momento asumiremos que los patrones representados por $y_i = +1$ y los representados por $y_i = -1$ son “linealmente separables”. Es decir que existe la función f que satisface la ecuación 1.8. Entonces, si son “linealmente separables” el hiperplano que separa ambas clases o categorías esta dado por

$$\boldsymbol{\beta}^T \mathbf{x} + b = 0 \quad (1.9)$$

lo que es equivalente a escribir

$$\begin{aligned} \boldsymbol{\beta}^T \mathbf{x}_i + b &\geq 0 \quad \text{para } y_i = +1 \\ \boldsymbol{\beta}^T \mathbf{x}_i + b &< 0 \quad \text{para } y_i = -1 \end{aligned} \quad (1.10)$$

Para un vector de parámetros $\boldsymbol{\beta}$ y sesgo “ b ” dados, la separación entre el hiperplano definido por 1.9 y el dato más cercano se denomina “Margen de separación ρ ”. Denominaremos a hiperplano óptimo aquella solución para la cual ρ es máximo.

El objetivo ahora es encontrar el vector de parámetros $\hat{\boldsymbol{\beta}}$ que defina ese hiperplano óptimo, dado el conjunto de entrenamiento $F = \{(\mathbf{x}_i, y_i)\}_{i=1}^N$. En particular llamaremos “vectores de soporte” a aquellos puntos $(\mathbf{x}_i, y_i)$ para los cuales se satisface la igualdad de la siguiente ecuación:

$$\begin{aligned} \boldsymbol{\beta}^T \mathbf{x}_i + b &\geq 1 \quad \text{para } y_i = +1 \\ \boldsymbol{\beta}^T \mathbf{x}_i + b &\leq -1 \quad \text{para } y_i = -1 \end{aligned} \quad (1.11)$$

En este contexto, aquellos vectores que están exactamente a la distancia ρ del hiperplano satisfacen la igualdad en la 1.11. Si consideramos el vector de soporte $\mathbf{x}^{(s)}$ para el cual tenemos que $y^{(s)} = \mp 1$, por la 1.11 tenemos que

$$g(\mathbf{x}^{(s)}) = \boldsymbol{\beta}_0^T \mathbf{x}^{(s)} \mp b = \mp 1 \quad \text{para } y^{(s)} = \mp 1.$$

Se puede demostrar que el margen ρ queda definido como $\rho = \frac{g(\mathbf{x}^{(s)})}{\|\boldsymbol{\beta}_0\|}$, con lo cual quedan determinadas las condiciones necesarias para la búsqueda del vector de parámetros óptimo de la siguiente manera:

Dado el conjunto de entrenamiento $F = \{(\mathbf{x}_i, y_i)\}_{i=1}^N$, encontrar los valores óptimos del vector de parámetros $\boldsymbol{\beta}$ y sesgo “ b ” tal que satisfagan las siguientes condiciones

$$y_i(\boldsymbol{\beta}^T \mathbf{x}_i + b) \geq 1 \quad \text{para } i=1, 2, \dots, N$$

y con el vector de parámetros β minimizando la función de costo o pérdida

$$\varphi(\beta) = \frac{1}{2} \beta^T \beta.$$

de las condiciones anteriores podemos observar las siguientes características:

- La función de costo $\varphi(\beta)$ es convexa en β .
- Las condiciones son lineales en β .

Este problema de optimización puede resolverse mediante la aplicación de los multiplicadores de Lagrange, lo cual implica reescribir las condiciones de la siguiente manera:

$$J(\beta, b, \alpha) = \frac{1}{2} \beta^T \beta - \sum_{i=1}^N \alpha_i [y_i (\beta^T x_i + b) - 1], \quad \alpha \geq 0 \quad (1.12)$$

La función J tiene forma de silla de montar que debe ser minimizada con respecto a β y b y maximizada con respecto a α . La solución de esto lleva a las siguientes ecuaciones que representan los valores óptimos de los parámetros:

$$\beta = \sum_{i=1}^N \alpha_i y_i x_i \quad (1.13)$$

La solución para β parece implicar la sumatoria sobre todos los vectores existentes en el conjunto de entrenamiento, sin embargo solamente los coeficientes de Lagrange α que satisfacen $\alpha_i [y_i (\beta^T x_i + b) - 1] = 0$ son los únicos que toman valores distintos de cero, reduciendo al 1.13 a

$$\beta = \sum_{i \in VS} \alpha_i y_i x_i$$

donde VS denota el subconjunto de F que satisface la igualdad de la 1.11.

Esta representación, conocida como "primal" tiene también una representación dual de la siguiente forma:

Dado el conjunto de entrenamiento $F = \{(x_i, y_i)\}_{i=1}^N$, encuentre los multiplicadores de Lagrange $\{\alpha_i\}_{i=1}^N$ que maximizan la función objetivo

$$Q(\alpha) = \sum_{i=1}^N \alpha_i - \frac{1}{2} \sum_{i=1}^N \sum_{j=1}^N \alpha_i \alpha_j y_i y_j x_i^T x_j$$

sujeto a las condiciones

1. $\sum_{i=1}^N \alpha_i y_i = 0$
2. $\alpha_i \geq 0$ para $i=1, 2, \dots, N$

Ahora bien, ¿Cómo se relaciona las Máquinas de Soporte Vectorial con lo expuesto sobre la teoría de la minimización del riesgo empírico?

Mediante el siguiente Teorema [Vapnik, 1995, 1998]

Sea D el diámetro de la "bola" más pequeña que contiene a todos los vectores de entrada $x_1, x_2, \dots, x_N$. El conjunto de hiperplanos óptimo descrito por la ecuación $\beta_0^T x + b_0 = 0$ tiene una dimensión VC "h" restringida superiormente por

$$h \leq \min \left\{ \left\lceil \frac{D^2}{\rho^2} \right\rceil, m_0 \right\} + 1$$

donde $[\cdot]$ refiere al valor entero mínimo que es mayor o igual al valor encerrado entre los corchetes, ρ es el margen de separación que es igual a $1/\|\beta_0\|$, y m_0 es la dimensión del espacio de entradas.

Ahora bien, hasta el momento solo hemos abordado la situación "ideal" de que los datos sean linealmente separables, pero en la realidad eso no sucede muy a menudo y en general existe cierto grado de solapamiento entre las distintas clases a clasificar, más aún cuando la dimensión del espacio de entradas es grande. Sin embargo, cuando este es el caso, aun podemos encontrar un hiperplano óptimo que minimice la probabilidad del error de clasificación promedio sobre el conjunto de entrenamiento.

En este contexto decimos que el margen de separación entre las clases es "blando" si existen o permitimos puntos $(\mathbf{x}_i, y_i)$ que violen la condición $y_i(\beta^T \mathbf{x}_i + b) \geq 1$. La violación a esta condición puede devenir de dos situaciones:

- El punto $(\mathbf{x}_i, y_i)$ cae dentro de la región de separación
- El punto $(\mathbf{x}_i, y_i)$ cae del lado opuesto al que debería según el valor de y_i , es decir del lado equivocado con respecto a la superficie de decisión.

Para formalizar el tratamiento de estas dos situaciones, introducimos un nuevo conjunto de variables (variables de holgura o amplitud) $\{\xi_i\}_{i=1}^N$ en la definición del hiperplano de separación según:

$$y_i(\beta^T \mathbf{x}_i + b) \geq 1 - \xi_i$$

Las variables de holgura ξ_i miden la desviación de cada punto con respecto a la condición ideal. Cuando $0 \leq \xi_i \leq 1$ el dato cae dentro de la región de separación (dentro del margen) pero del lado correcto con respecto al hiperplano. Mientras que si $\xi_i > 1$ el punto cae en el lugar equivocado del plano de decisión. Para encontrar ahora el vector de parámetros que optimiza el margen, reescribimos el funcional de costo como sigue:

$$\varphi(\beta, \xi) = \frac{1}{2} \beta^T \beta + C \sum_{i=1}^N \xi_i$$

La constante $C > 0$ cumple el rol de parámetro de "regularización", que controla el compromiso entre la complejidad de la MSV y el número de vectores de soporte no-separables. Este parámetro C es usualmente seleccionado por el usuario.

La forma dual del problema de optimización queda ahora expresada de la siguiente forma:

Dado el conjunto de entrenamiento $F = \{(\mathbf{x}_i, y_i)\}_{i=1}^N$, encuentre los multiplicadores de Lagrange $\{\alpha_i\}_{i=1}^N$ que maximizan la función objetivo

$$Q(\alpha) = \sum_{i=1}^N \alpha_i - \frac{1}{2} \sum_{i=1}^N \sum_{j=1}^N \alpha_i \alpha_j y_i y_j \mathbf{x}_i^T \mathbf{x}_j$$

sujeito a las condiciones

1. $\sum_{i=1}^N \alpha_i y_i = 0$
2. $0 \leq \alpha_i \leq C$ para $i=1, 2, \dots, N$

donde C es un parámetro positivo especificado por el usuario.

Como se puede ver, en la formulación dual no aparecen las “variables de holgura”, siendo la solución prácticamente la misma que en el caso de margen rígido, excepto por una pequeña pero importante diferencia, ahora los α_s tienen una cota superior C .

El enlace entre la solución primal y dual se mantiene como

$$\boldsymbol{\beta} = \sum_{i \in VS} \alpha_i y_i \mathbf{x}_i$$

Solo que ahora tenemos que si $\alpha_i < C$, entonces $\xi_i = 0 \Rightarrow \mathbf{x}_i$ está del lado correcto del hiperplano, en cambio si $\alpha_i = C$, entonces $\xi_i > 0 \Rightarrow \mathbf{x}_i$ está dentro del margen si $0 < \xi_i < 1$ o del lado incorrecto del hiperplano $\xi_i \geq 1$.

El principio de MRE en regresión

En este caso, la variable $y \in \mathcal{R}$. En el abordaje clásico e incluso en métodos basados en redes neuronales artificiales, la función de pérdida es la cuadrática i.e $(y - f(\mathbf{x}, \boldsymbol{\beta}))^2$. Sin embargo, este estimador de riesgo es muy sensible a valores infrecuentes ("outliers") como a distribuciones asimétricas. Para abordar este problema, se requieren estimadores robustos e insensibles a pequeños cambios en el modelo.

Como una extensión del modelo basado en SVM que se ajusta a los criterios del MRE, Vapnik propuso en 1995 la siguiente función de pérdida:

$$L_\epsilon^p(f(\mathbf{x}, \boldsymbol{\beta}), y) = \begin{cases} |f(\mathbf{x}, \boldsymbol{\beta}) - y|^p - \epsilon, & \text{para } |f(\mathbf{x}, \boldsymbol{\beta}) - y|^p \geq \epsilon \\ 0 & \text{en cc} \end{cases} \text{ with } p = 1, 2$$

donde ϵ es el parámetro de tolerancia. Esta función de pérdida L_ϵ^p se denomina ϵ - *insensible* de norma p .

Ahora podemos suponer que estamos ante el siguiente problema:

$$y = f(\mathbf{x}) + v$$

donde la función $f(\mathbf{x})$ (posiblemente no lineal) está definida por la esperanza condicional $E(Y|\mathbf{x})$ donde Y es una variable aleatoria con realizaciones " y ". El ruido aditivo " v " es estadísticamente independiente de " $\mathbf{x}$ " y tanto la función $f()$ como " v " son desconocidos. Todo lo que tenemos disponible es el conjunto de datos de entrenamiento $F = \{(\mathbf{x}_i, y_i)\}_{i=1}^N$. Otra vez aquí, el problema consiste en minimizar el riesgo empírico

$$R_{emp} = \frac{1}{N} \sum_{i=1}^N L_\epsilon(y_i, \mathbf{x}_i)$$

Y para el caso de L_ϵ^1 podemos formular el problema de la siguiente manera:

$$\begin{aligned} y_i - \boldsymbol{\beta}^T \mathbf{x}_i &\leq \epsilon - \xi_i \\ \boldsymbol{\beta}^T \mathbf{x}_i - y_i &\leq \epsilon - \xi'_i \end{aligned}$$

donde $\{\xi_i \geq 0\}_{i=1}^N$ son variables de holgura. Este problema de optimización es equivalente a minimizar el siguiente funcional de costo

$$\Phi(\boldsymbol{\beta}, \boldsymbol{\xi}, \boldsymbol{\xi}') = C \left(\sum_{i=1}^N (\xi_i + \xi'_i) \right) + \frac{1}{2} \boldsymbol{\beta}^T \boldsymbol{\beta}$$

Sujeto a

$$\begin{aligned} ((\boldsymbol{\beta} \cdot \mathbf{x}_i) + b) - y_i &\leq \epsilon + \xi_i \\ y_i - ((\boldsymbol{\beta} \cdot \mathbf{x}_i) + b) &\geq \epsilon + \xi'_i \end{aligned}$$

$$\xi_i, \xi'_i \geq 0, i = 1, 2, \dots, N$$

donde la constante C es un parámetro especificado por el usuario.

Mediante la aplicación de los multiplicadores de Lagrange, este problema de optimización se puede escribir de la siguiente manera:

$$B(\alpha) = \sum_{i=1}^N y_i \cdot \alpha_i - \varepsilon \sum_{i=1}^N |\alpha_i| - \frac{1}{2} \sum_{i,j=1}^N \alpha_i \alpha_j x_i x_j$$

Sujeto a

$$\sum_i \alpha_i = 0, -C \leq \alpha_i \leq C, i = 1..N$$

Como en el caso anterior tenemos que solo unos pocos coeficientes de Lagrange son distintos de cero y definen a los vectores de soporte, por lo que podemos escribir el vector

$$\beta = \sum_{i \in SV} (\alpha_i) x_i$$

Luego la ecuación de predicción queda como

$$y_j = \beta x_j = \sum_{i \in SV} (\alpha_i) x_i x_j$$

Referencias:

Haykin, S. Neural Networks and Learning Machines 3ed. Prentice Hall. 2008
Vapnik V. The nature of Statistical Learning Theory. 2ed. Springer 1999.