

IN550 Machine Learning

Debugging dell'apprendimento supervisionato

Vincenzo Bonifaci

Esempio di scenario

- Volete addestrare un filtro spam
- Avete scelto un insieme di 100 parole da utilizzare come feature (invece di utilizzare 50000+ parole della lingua)
- La regressione logistica regolarizzata, implementata tramite discesa del gradiente, ottiene un errore di test del 20%, che per la vostra applicazione è **inaccettabile**

Come procedere?

Messa a punto del processo di apprendimento

- Approccio possibile: provare a migliorare il processo in diversi modi
 - Ottenendo più esempi di training
 - Provando un insieme di feature più piccolo
 - Provando un insieme di feature più grande
 - Usando feature diverse: intestazione mail vs. corpo mail
 - Aumentando il numero di iterazioni di Gradient Descent
 - Usando un metodo di ottimizzazione del secondo ordine
 - Usando un diverso valore del coefficiente di regolarizzazione λ
 - Usando una Support Vector Machine anziché la regressione logistica
- Provare tutto richiede molto tempo, e rischia di diventare una questione di fortuna

Diagnosi di bias e varianza

Approccio preferibile:

- Diagnosi del problema incontrato dall'apprendimento
- Correzione del problema

Nel nostro scenario, l'errore di test (20%) è troppo alto

Supponiamo che sospettiate che il problema sia uno tra i seguenti due:

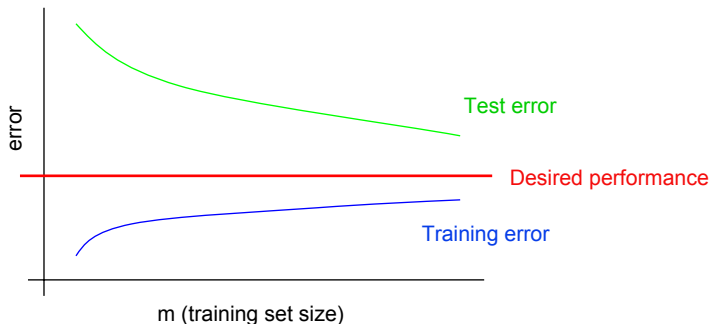
- Overfitting (**varianza** alta)
- Troppe poche feature per una corretta classificazione dello spam (**bias** alto)

Diagnosi:

- Varianza alta: l'errore di training sarà **sostanzialmente inferiore** all'errore di test
- Bias alto: l'errore di training sarà **simile** a quello di test

Bias vs. varianza e curve di apprendimento

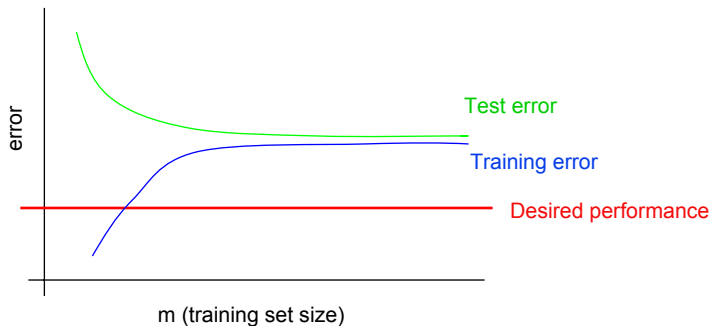
Curva di apprendimento tipica quando la **varianza** è alta:



- L'errore di test continua a decrescere con m crescente. Ciò suggerisce che un training set più grande possa essere utile
- Grande gap tra errore di test ed errore di training

Bias vs. varianza e curve di apprendimento

Curva di apprendimento tipica quando il **bias** è alto:



- Già l'errore di training è troppo alto!
- Piccolo gap tra errore di test ed errore di training

La diagnosi ci dice come intervenire

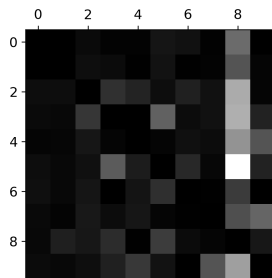
Nel nostro scenario, si può intervenire:

- Ottenendo più esempi di training Diminuisce la varianza
- Provando un insieme di feature più piccolo Diminuisce la varianza
- Provando un insieme di feature più grande Diminuisce il bias
- Usando l'intestazione della mail Diminuisce il bias
- Aumentando il numero di iterazioni di Gradient Descent Corregge l'algoritmo di ottimizzazione
- Usando un metodo del secondo ordine Corregge l'algoritmo di ottimizzazione
- Usando un diverso valore del coefficiente di regolarizzazione λ
 - Aumentare λ diminuisce la varianza
 - Diminuire λ diminuisce il bias
- Usando una Support Vector Machine Corregge l'obiettivo dell'ottimizzazione

Analisi degli errori

Esempio: classificazione cifre MNIST

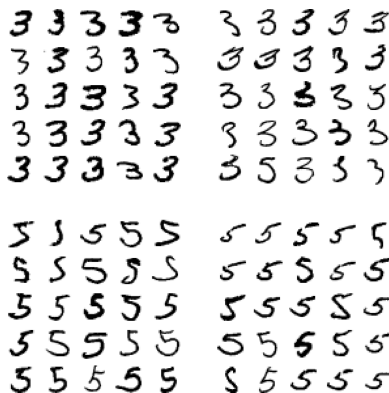
Analizziamo la **matrice di confusione** per capire quali classi vengono confuse con quali altre (**nota**: la diagonale è stata azzerata)



- Molte cifre (ad es., tanti 5) vengono facilmente confuse con degli 8
- I veri 8 vengono per lo più classificati correttamente
- Sembra utile aggiungere esempi che sembrano 8 ma non lo sono

Analisi degli errori

Analizziamo **errori individuali** per capire perché il classificatore fallisce



- Il classificatore sembra troppo sensibile a traslazione e rotazione
- Ciò suggerisce di centrare e “raddrizzare” le immagini

- Il tempo investito nella diagnosi dell'algoritmo di apprendimento è tempo ben investito
- La diagnosi di bias e varianza è molto utile e può essere semplice
- L'analisi degli errori individuali può fornire altre informazioni utili