

Deep Learning

Revisão de *Machine Learning*

Lucas Silveira Kupssinskü

Machine Learning Theory and Applications Laboratory
PUCRS

agosto de 2026

Overview

1. Por que precisamos de Machine Learning?

2. Conceitos Básicos de *Machine Learning*

3. Modelos Lineares

3.1 Regressão Linear

3.2 Regressão Logística

Visão Geral

1. Por que precisamos de Machine Learning?

2. Conceitos Básicos de *Machine Learning*

3. Modelos Lineares

3.1 Regressão Linear

3.2 Regressão Logística

A importância de Machine Learning

- “A breakthrough in machine learning would be worth ten Microsofts” (Bill Gates)
- “Machine learning is the next Internet” (Tony Tether, Director, DARPA)
- “Machine learning is going to result in a real revolution” (Greg Papadopoulos, Former CTO, Sun)
- “Machine learning is today’s discontinuity” (Jerry Yang, Founder, Yahoo)
- “Machine learning is the new electricity” (Andrew NG, Stanford Professor and Coursera founder)

Software sem Machine Learning

É uma sequência de instruções que:

- Recebe uma entrada
- Realiza um processamento
- Retorna uma saída

Software sem Machine Learning

Exemplo

1. Se você ou o seu oponente tiverem dois em sequência, jogue no espaço restante
2. Caso contrário, se existir um movimento que crie duas sequencias de dois, faça essa jogada
3. Caso contrário, se o centro estiver vazio, jogue no centro
4. Caso contrário, se o seu oponente jogou em um canto, jogue no canto oposto
5. Caso contrário, se houver um canto livre jogue nesse canto
6. Caso contrário, jogue em uma posição vazia qualquer

Algumas vezes não sabemos como resolver

- Complexidade de Tempo
- Complexidade de Espaço
- Complexidade “Humana”



E no Aprendizado de Máquina?

- Mesmo quando não é possível criar um algoritmo para resolver o problema diretamente, pode ser possível criar algoritmos que *encontram padrões*
- Usamos raciocínio *indutivo* para resolver os problemas

Definição de *Machine Learning*

“A computer program is said to learn from experience **E** with respect to some class of tasks **T** and performance measure **P**, if its performance at tasks in **T**, as measured by **P**, improves with experience **E**” (Mitchell, 1997)

Tipos de Raciocínio

- **Dedutivo**

- Assumindo algumas premissas, a conclusão obrigatoriamente segue

- **Não-dedutivo**

- **Abdução**

- “Raciocínio do Detetive”

- **Indução**

- Dadas observações (limitadas em escopo e quantidade), podemos chegar a conclusões plausíveis

Exemplo de Dedução

- Premissa 1: Todos os humanos são mortais
- Premissa 2: Sócrates é humano
- Conclusão: Sócrates é mortal

Tipos de Raciocínio

- **Dedutivo**

- Assumindo algumas premissas, a conclusão obrigatoriamente segue

- **Não-dedutivo**

- **Abdução**

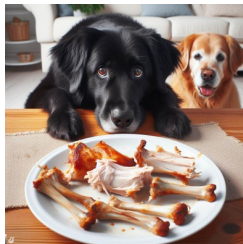
- “Raciocínio do Detetive”

- **Indução**

- Dadas observações (limitadas em escopo e quantidade), podemos chegar a conclusões plausíveis

Exemplo de Abdução

- Tem restos de frango assado pelo chão
- O cachorro está com restos de frango no pelo, especialmente ao redor da boca
- Logo, o cachorro comeu o frango assado



Tipos de Raciocínio

- **Dedutivo**

- Assumindo algumas premissas, a conclusão obrigatoriamente segue

- **Não-dedutivo**

- **Abdução**

- “Raciocínio do Detetive”

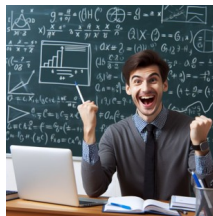
- **Indução**

- Dadas observações (limitadas em escopo e quantidade), podemos chegar a conclusões plausíveis

Exemplo de Indução

Qual o próximo número na sequência abaixo?

1, 3, 5, 7, x



Tipos de Raciocínio

- **Dedutivo**

- Assumindo algumas premissas, a conclusão obrigatoriamente segue

- **Não-dedutivo**

- **Abdução**

- “Raciocínio do Detetive”

- **Indução**

- Dadas observações (limitadas em escopo e quantidade), podemos chegar a conclusões plausíveis

Exemplo de Indução

Qual o próximo número na sequência abaixo?

1, 3, 5, 7, **217341**

$$f(x) = \frac{18111}{2}x^4 - 90555x^3 + \frac{633885}{2}x^2 - 452773x + 217331$$



1. Por que precisamos de Machine Learning?

2. Conceitos Básicos de *Machine Learning*

3. Modelos Lineares

3.1 Regressão Linear

3.2 Regressão Logística

Desenvolvimento de uma aplicação com AM

O projeto de um sistema que usa *Machine Learning* pode ser descrito pela figura abaixo¹

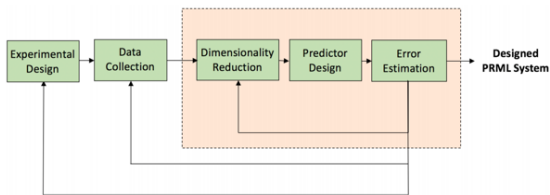


Figura: Representação simplificada do processo de *Machine Learning*. Adaptado de (Braga-Neto, 2020)

¹Fluxos diferentes existem na literatura de *Knowledge Discovery in Databases* e em Mineração de Dados.

Desenvolvimento de uma aplicação com AM

Experimental Design

- Modelar o problema como uma tarefa de *Machine Learning*
- Como você definiria o problema de treinar um reconhecedor de rostos que é capaz de reconhecer 100 pessoas diferentes?
 - Como um problema de **classificação multiclasse** com 100 classes será muito difícil coletar uma base dados grande o bastante
 - Como um problema de **classificação binário** onde o sistema recebe um par de imagens é mais fácil coletar dados

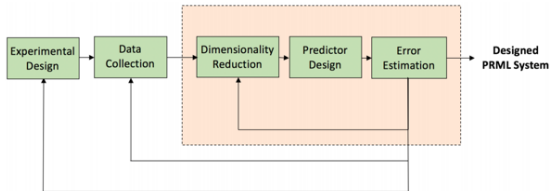


Figura: Representação simplificada do processo de *Machine Learning*. Adaptado de (Braga-Neto, 2020)

Desenvolvimento de uma aplicação com AM

Data Collection

- É possível coletar dados que sejam representativos do problema em questão?
- Existe algum dataset aberto que pode nos auxiliar?
- Quando a coleta de dados ou a anotação é cara, existem alternativas?

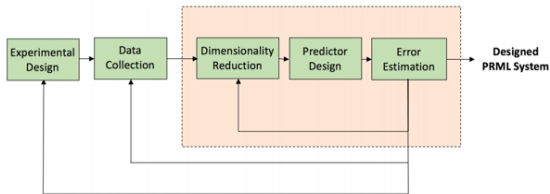


Figura: Representação simplificada do processo de *Machine Learning*. Adaptado de (Braga-Neto, 2020)

Desenvolvimento de uma aplicação com AM

Dimensionality Reduction

- Todos os atributos coletados são úteis para resolver a tarefa?
- Como selecionar apenas os atributos relacionados a variável alvo?
- Existem formas de criar novos atributos que sejam melhores preditores para a variável alvo?
- Posso resumir meus atributos em um espaço menor?

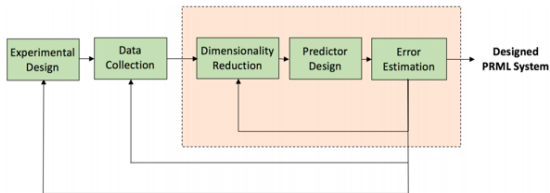


Figura: Representação simplificada do processo de *Machine Learning*. Adaptado de (Braga-Neto, 2020)

Desenvolvimento de uma aplicação com AM

Predictor Design

- Existe algum modelo cujo viés indutivo seja aderente ao meu problema?
- Qual será a métrica utilizada para avaliação do desempenho do modelo?
- Qual meu *budget* computacional?
- Qual a melhor hiperparametrização dos modelos selecionados?

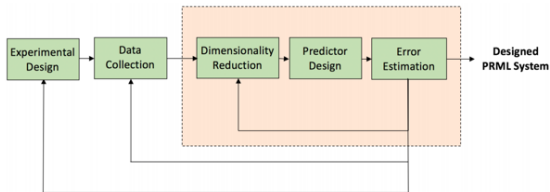


Figura: Representação simplificada do processo de *Machine Learning*. Adaptado de (Braga-Neto, 2020)

Desenvolvimento de uma aplicação com AM

Error Estimation

- Meu modelo treinado generaliza?
- Qual é o erro esperado do meu modelo para novos dados?
- Coletar mais dados pode melhorar meus resultados?

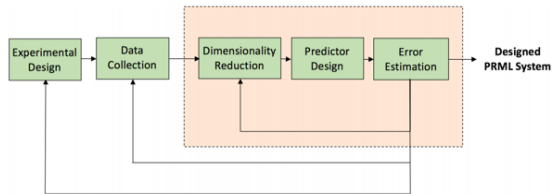


Figura: Representação simplificada do processo de *Machine Learning*. Adaptado de (Braga-Neto, 2020)

Tipos de Aprendizado

- **Supervisionado**

- Dados de treino possuem rótulos
- *Exemplo:* classificar fotos como gato ou cachorro a partir de 10 mil imagens rotuladas

- **Não Supervisionado**

- Dados de treino não possuem rótulos
- *Exemplo:* agrupar clientes em segmentos a partir do histórico de compras, sem rótulos

- **Semi Supervisionado**

- Alguns dados de treino possuem rótulo outros não
- *Exemplo:* detecção de fraude com poucas transações rotuladas e milhões sem rótulo

- **Aprendizado por Reforço**

- Agente recebe “recompensa” devido a uma série de ações
- *Exemplo:* AlphaGo aprende a jogar Go ganhando recompensa ao final da partida

- **Auto Supervisionado**

- Parte das entradas servem como “supervisão”
- *Exemplo:* pré-treino do BERT prevendo palavras mascaradas em uma frase

Notação

Considerando um dataset formado por exemplos de uma relação $(\mathbf{X}^{(i)}, \mathbf{y}^{(i)})$

- $\mathbf{X}$ é nossa matriz de atributos $N \times M$ onde N é o número de instâncias e M é o número de atributos indexada por $\mathbf{X}_m^{(n)}$
- $\mathbf{y}$ é nossa variável alvo $N \times 1$
- Objetivo é estimar $\mathbf{y}^{(i)}$ para uma instância desconhecida $\mathbf{X}^{(i)}$
 - Quando $\mathbf{y}^{(i)}$ é uma variável discreta: *Problema de Classificação*
 - Quando $\mathbf{y}^{(i)}$ é uma variável contínua: *Problema de Regressão*

Viés Indutivo

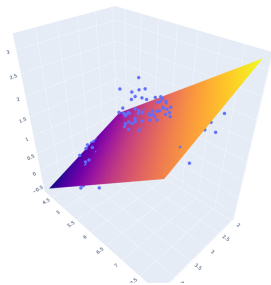


Figura: Regressão Linear

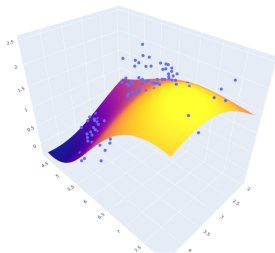


Figura: SVM

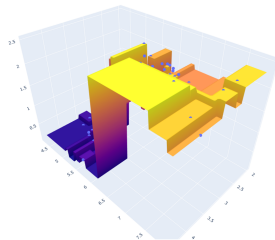


Figura: Árvore de Decisão

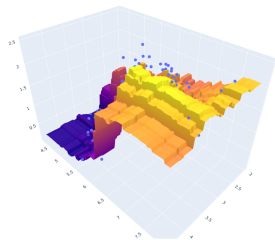


Figura: *Random Forest*

Efeito tesoura (*Scissor Effect*)

- Quando temos um conjunto de dados pequeno, regras de classificação mais simples tendem a obter erros de classificação menores
- Quando temos um conjunto de dados grande, regras de classificação mais complexas tendem a obter erros menores

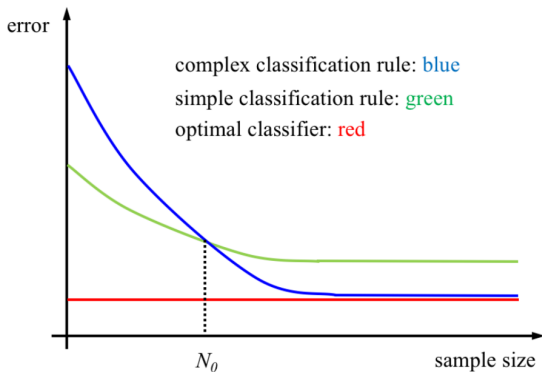
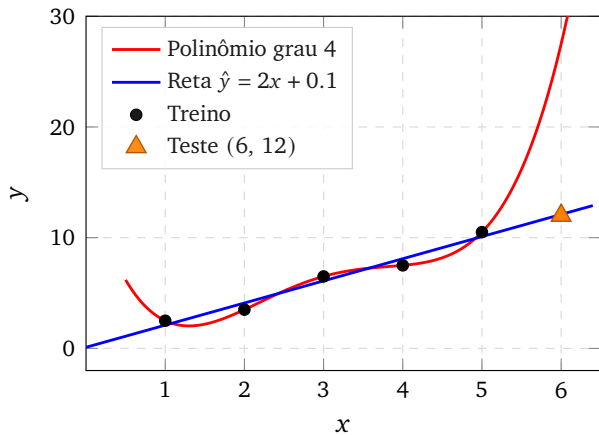


Figura: Adaptado de (Braga-Neto, 2020)

Efeito tesoura, exemplo numérico



Em $x = 6$: reta prediz 12.1 (acerta), polinômio prediz 27.5 (erra muito).

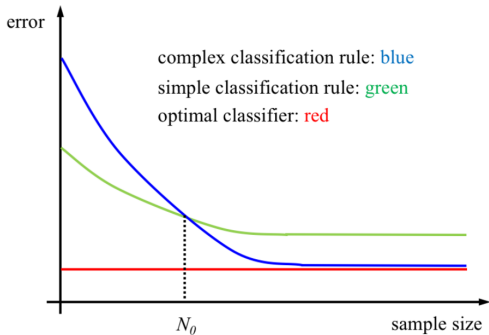


Figura: Adaptado de (Braga-Neto, 2020)

Trade-off Viés-Variância

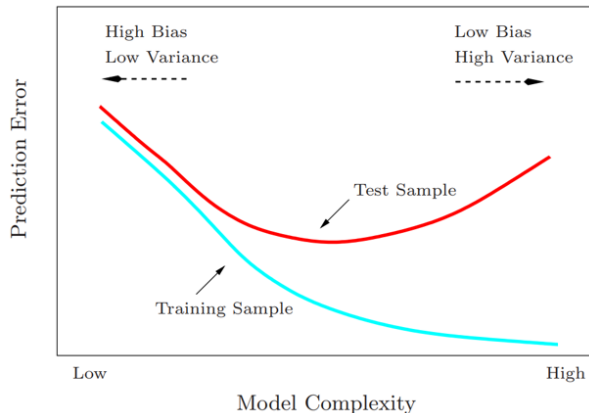


Figura: Adaptado de (Braga-Neto, 2020).

Decomposição

$$\text{Erro}_{\text{teste}}(x) = \text{Viés}^2(x) + \text{Variância}(x) + \sigma^2$$

- **Viés alto:** modelo simples demais, não captura o padrão.
- **Variância alta:** modelo sensível à amostra de treino.
- **Ruído (σ^2):** limite irreduzível.

Derivação completa em *Regularização* (Aula 08).

Visão Geral

1. Por que precisamos de Machine Learning?

2. Conceitos Básicos de *Machine Learning*

3. Modelos Lineares

3.1 Regressão Linear

3.2 Regressão Logística

Modelos Lineares

- Modelos lineares baseados em variações dos Método dos Mínimos Quadrados talvez sejam o tipo de modelo mais estudado por todas as áreas do conhecimento.
- Faremos uma breve revisão com viés de *Machine Learning*

Intuição

Queremos ajustar um *hiperplano*¹ aos nossos dados

$$\hat{y} = \mathbf{X}_1 \cdot \theta_1 + \theta_0$$

¹No $\mathbb{R}^2$ um hiperplano é equivalente a uma reta

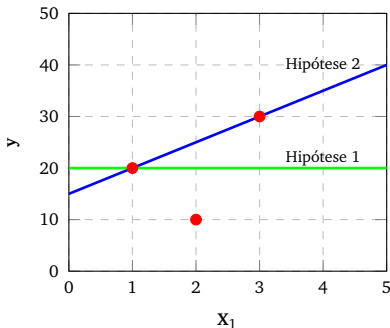
Intuição

Queremos ajustar um *hiperplano*¹ aos nossos dados

$$\hat{y} = \mathbf{X}_1 \cdot \theta_1 + \theta_0$$

X	y
1	20
2	10
3	30

Tabela:
Conjunto
de Dados



Hipótese 1

$$\theta_0 = 20, \theta_1 = 0$$

$$\hat{y} = 0X_1 + 20$$

Hipótese 2

$$\theta_0 = 15, \theta_1 = 5$$

$$\hat{y} = 5X_1 + 15$$

¹No $\mathbb{R}^2$ um hiperplano é equivalente a uma reta

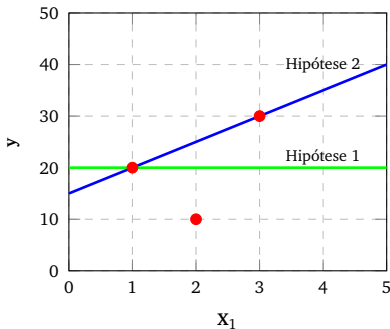
Intuição

Qual é a melhor hipótese?

$$J_{SSR} = \sum_{i=1}^N \left(\hat{y}^{(i)} - y^{(i)} \right)^2$$

X	y
1	20
2	10
3	30

Tabela:
Conjunto
de Dados



Hipótese 1

$$\theta_0 = 20, \theta_1 = 0$$

$$\hat{y} = 0x_1 + 20$$

Hipótese 2

$$\theta_0 = 15, \theta_1 = 5$$

$$\hat{y} = 5x_1 + 15$$

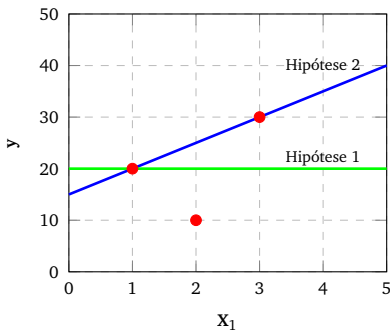
Intuição

Qual é a melhor hipótese?

$$J_{SSR} = \sum_{i=1}^N \left(\hat{y}^{(i)} - y^{(i)} \right)^2$$

X	y
1	20
2	10
3	30

Tabela:
Conjunto
de Dados



✓ Hipótese 1

$$\theta_0 = 20, \theta_1 = 0$$

$$\hat{y} = 0x_1 + 20$$

$$J_{SSR} = (20 - 20)^2 + (10 - 20)^2 + (30 - 20)^2 = 200$$

✗ Hipótese 2

$$\theta_0 = 15, \theta_1 = 5$$

$$\hat{y} = 5x_1 + 15$$

$$J_{SSR} = (20 - 20)^2 + (10 - 25)^2 + (30 - 30)^2 = 225$$

Treinamento

Como encontrar a melhor hipótese?

- $\operatorname{argmin}_{\theta_0, \theta_1} [J_{\text{SSR}}]$
- $J_{\text{SSR}} = \sum_{i=1}^N (\hat{\mathbf{y}}^{(i)} - \mathbf{y}^{(i)})^2 = (\hat{\mathbf{y}} - \mathbf{y})^\top (\hat{\mathbf{y}} - \mathbf{y})$
- $\hat{\mathbf{y}} = \mathbf{X}_1 \cdot \theta_1 + \theta_0$

$\mathbf{X}_1$	$\mathbf{y}$		
1	20	θ_0	θ_1
2	10		
3	30		

Notation Trick

- Colocamos o parâmetro livre junto do vetor θ fazendo um *truque de notação*
- Imagine que existe um novo atributo cujo valor é 1 para todas as instâncias
- Assim todos os parâmetros treináveis podem ser tratados da mesma forma.
- $\hat{y} = X\theta$

X_0	X_1		
1	1	y	θ
1	2	20	θ_0
1	3	10	θ_1
		30	
$N \times M$		$N \times 1$	$M \times 1$

Treinamento

Como encontrar a melhor hipótese?

- $\operatorname{argmin}_{\theta_0, \theta_1} [J_{\text{SSR}}]$
- $J_{\text{SSR}} = \sum_{i=1}^N (\hat{\mathbf{y}}^{(i)} - \mathbf{y}^{(i)})^2 = (\hat{\mathbf{y}} - \mathbf{y})^\top (\hat{\mathbf{y}} - \mathbf{y})$
- $\hat{\mathbf{y}} = \mathbf{X}\boldsymbol{\theta}$

$\mathbf{X}_0$	$\mathbf{X}_1$	$\mathbf{y}$	$\boldsymbol{\theta}$
1	1	20	θ_0
1	2	10	θ_1
1	3	30	
$N \times M$		$N \times 1$	$M \times 1$

Sabendo que..

- A soma dos quadrados dos resíduos é uma *função de custo convexa* para o problema da regressão linear
- Logo, *possui apenas um mínimo*
- Podemos utilizar o gradiente da função de custo em um processo iterativo de otimização

Treinamento

Vamos derivar J_{SSR} em relação a θ e igualar a 0

$$\begin{aligned}\frac{\partial J_{\text{SSR}}}{\partial \theta} &= \frac{\partial ((\hat{\mathbf{y}} - \mathbf{y})^\top (\hat{\mathbf{y}} - \mathbf{y}))}{\partial \theta} \\&= \frac{\partial (\hat{\mathbf{y}}^\top \hat{\mathbf{y}} - \hat{\mathbf{y}}^\top \mathbf{y} - \mathbf{y}^\top \hat{\mathbf{y}} + \mathbf{y}^\top \mathbf{y})}{\partial \theta} \\&= \frac{\partial (\hat{\mathbf{y}}^\top \hat{\mathbf{y}} - 2\hat{\mathbf{y}}^\top \mathbf{y} + \mathbf{y}^\top \mathbf{y})}{\partial \theta} \\&= \frac{\partial (\boldsymbol{\theta}^\top \mathbf{X}^\top \mathbf{X} \boldsymbol{\theta} - 2\boldsymbol{\theta}^\top \mathbf{X}^\top \mathbf{y} + \mathbf{y}^\top \mathbf{y})}{\partial \theta} \\&= \frac{\partial (\boldsymbol{\theta}^\top \mathbf{X}^\top \mathbf{X} \boldsymbol{\theta})}{\partial \theta} - \frac{\partial (2\boldsymbol{\theta}^\top \mathbf{X}^\top \mathbf{y})}{\partial \theta} \\&= 2\mathbf{X}^\top \mathbf{X} \boldsymbol{\theta} - 2\mathbf{X}^\top \mathbf{y} \\&= 2\mathbf{X}^\top (\mathbf{X} \boldsymbol{\theta} - \mathbf{y})\end{aligned}$$

Treinamento

Vamos derivar J_{SSR} em relação a θ e igualar a 0

$$2\mathbf{X}^\top(\mathbf{X}\theta - \mathbf{y}) = 0$$

$$\mathbf{X}^\top(\mathbf{X}\theta - \mathbf{y}) = 0$$

$$\mathbf{X}^\top\mathbf{X}\theta - \mathbf{X}^\top\mathbf{y} = 0$$

$$\mathbf{X}^\top\mathbf{X}\theta = \mathbf{X}^\top\mathbf{y}$$

$$\theta = (\mathbf{X}^\top\mathbf{X})^{-1}\mathbf{X}^\top\mathbf{y}$$

Solução analítica para θ

$$\theta = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y}$$

$$\left(\begin{array}{|c|c|c|} \hline 1 & 1 & 1 \\ \hline 1 & 2 & 3 \\ \hline \end{array} \times \begin{array}{|c|c|} \hline 1 & 1 \\ \hline 1 & 2 \\ \hline 1 & 3 \\ \hline \end{array} \right)^{-1} \times \begin{array}{|c|c|c|} \hline 1 & 1 & 1 \\ \hline 1 & 2 & 3 \\ \hline \end{array} \times \begin{array}{|c|} \hline 20 \\ \hline 10 \\ \hline 30 \\ \hline \end{array} = \theta_{M \times 1}$$

$M \times N$ $N \times M$ $M \times N$ $N \times 1$

Solução analítica para θ

$$\theta = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y}$$

$$\left(\begin{array}{|c|c|} \hline 3 & 6 \\ \hline 6 & 14 \\ \hline \end{array} \right)_{M \times M}^{-1} \times \begin{array}{|c|c|c|} \hline 1 & 1 & 1 \\ \hline 1 & 2 & 3 \\ \hline \end{array}_{M \times N} \times \begin{array}{|c|} \hline 20 \\ \hline 10 \\ \hline 30 \\ \hline \end{array}_{N \times 1} = \theta_{M \times 1}$$

Solução analítica para θ

$$\theta = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y}$$

$$\begin{array}{|c|c|} \hline 2.33 & -1 \\ \hline -1 & 0.5 \\ \hline \end{array} \underset{M \times M}{\times} \begin{array}{|c|c|c|} \hline 1 & 1 & 1 \\ \hline 1 & 2 & 3 \\ \hline \end{array} \underset{M \times N}{\times} \begin{array}{|c|} \hline 20 \\ \hline 10 \\ \hline 30 \\ \hline \end{array} \underset{N \times 1}{\times} = \underset{M \times 1}{\theta}$$

Solução analítica para θ

$$\theta = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y}$$

$$\begin{array}{|c|c|c|} \hline 1.33 & 0.333 & -0.666 \\ \hline -0.5 & 0 & 0.5 \\ \hline \end{array} \underset{M \times N}{\times} \begin{array}{|c|} \hline 20 \\ \hline 10 \\ \hline 30 \\ \hline \end{array} \underset{N \times 1}{=} \underset{M \times 1}{\theta}$$

Solução analítica para θ

$$\theta = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y}$$

$$\begin{array}{|c|} \hline 10 \\ \hline 5 \\ \hline \end{array} = \underset{M \times 1}{\theta}$$

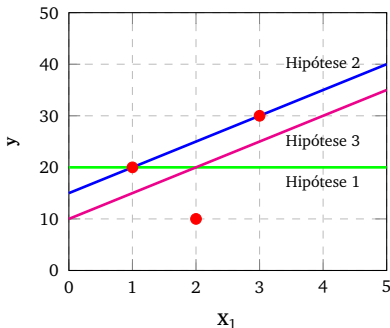
Sobre a solução analítica para θ

- Esse é o vetor de pesos que minimiza a soma dos quadrados dos resíduos (J_{SSR}).
- Repare que existe uma operação de inversão da matriz $\mathbf{X}^\top \mathbf{X}$. Essa operação gera um custo computacional $O(M^3)$ em relação à quantidade de atributos do problema
 - Esse custo pode inviabilizar a solução analítica

Qual é a melhor hipótese?

X	y
1	20
2	10
3	30

Tabela:
Conjunto
de Dados



❌ Hipótese 1

$$\theta_0 = 20, \theta_1 = 0, \hat{y} = 0X_1 + 20$$

$$J_{SSR} = (20 - 20)^2 + (10 - 20)^2 + (30 - 20)^2 = 200$$

❌ Hipótese 2

$$\theta_0 = 15, \theta_1 = 5, \hat{y} = 5X_1 + 15$$

$$J_{SSR} = (20 - 20)^2 + (10 - 25)^2 + (30 - 30)^2 = 225$$

✅ Hipótese 3

$$\theta_0 = 10, \theta_1 = 5, \hat{y} = 5X_1 + 10$$

$$J_{SSR} = (20 - 15)^2 + (10 - 20)^2 + (30 - 25)^2 = 150$$

Descida de Gradiente

- Algoritmo Iterativo
- $\theta_{t+1} = \theta_t - \eta \nabla_J$
- θ_0 inicializado com valores aleatórios
- $\eta \in (0, 1]$ taxa de aprendizado
- Valores típicos para η são 10^{-3} , 10^{-4} ou 10^{-5}

Descida de Gradiente

- Ao lado temos um algoritmo para *Batch Gradient Descent*
- Se selecionarmos apenas uma instância $\mathbf{X}^{(i)}$ para estimar o valor do gradiente teremos o *Stochastic Gradient Descent*
- A diferença entre J_{SSR} de diferentes iterações pode ser usado como condição de convergência

Require: $\mathbf{X}_{n \times m}$ matriz de instâncias

Require: $\mathbf{y}_{n \times 1}$ variável alvo

Require: $\eta \in (0, 1]$ taxa de aprendizado

```
1: procedure GRADIENT DESCENT
2:    $\boldsymbol{\theta} \sim \mathcal{N}(0, 1)$ 
3:   repeat
4:      $\boldsymbol{\theta} \leftarrow \boldsymbol{\theta} - \eta \mathbf{X}^\top (\mathbf{X}\boldsymbol{\theta} - \mathbf{y})$ 
5:      $\hat{\mathbf{y}} \leftarrow \mathbf{X}\boldsymbol{\theta}$ 
6:      $J_{SSR} \leftarrow (\mathbf{y} - \hat{\mathbf{y}})^\top (\mathbf{y} - \hat{\mathbf{y}})$ 
7:   until Condição de convergência
8:   return  $\boldsymbol{\theta}$ 
9: end procedure
```

Descida de Gradiente

Descida de Gradiente	Equação Normal
$\theta_{t+1} = \theta_t - \eta \mathbf{X}^\top (\mathbf{X}\theta_t - \mathbf{y})$	$\theta = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y}$
solução iterativa	solução analítica
Possui o hiperparâmetro η	Sem hiperparâmetros
arbitrariamente próxima ao melhor θ	encontra exatamente o melhor θ
$O(i \cdot N \cdot M)$ em implementações vetoriais	$O(M^3)$

Como aplicar Regressão Linear para Classificação?

A *imagem* da Regressão Linear é incompatível com uma classificação

- $p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)}) \in [0, 1]$
- $\mathbf{X}^{(i)} \boldsymbol{\theta} \in (-\infty, +\infty)$

Precisamos encontrar algum valor relacionado a probabilidade de pertencimento a uma classe $p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)})$ que tenha um contradomínio similar a Regressão Linear $\mathbf{X}^{(i)} \boldsymbol{\theta}$

Como aplicar Regressão Linear para Classificação?

- $p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)}) \in [0, 1]$: probabilidade de pertencimento a classe c
- $\frac{p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)})}{1 - p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)})} \in [0, +\infty)$: chance de pertencimento a classe c
- $\ln \left(\frac{p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)})}{1 - p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)})} \right) \in (-\infty, +\infty)$: log da chance de pertencimento a classe c :

Como aplicar Regressão Linear para Classificação?

- $p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)}) \in [0, 1]$: probabilidade de pertencimento a classe c
- $\frac{p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)})}{1 - p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)})} \in [0, +\infty)$: chance de pertencimento a classe c
- $\ln \left(\frac{p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)})}{1 - p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)})} \right) \in (-\infty, +\infty)$: log da chance de pertencimento a classe c :

Como aplicar Regressão Linear para Classificação?

- $p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)}) \in [0, 1]$: probabilidade de pertencimento a classe c
- $\frac{p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)})}{1 - p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)})} \in [0, +\infty)$: chance de pertencimento a classe c
- $\ln \left(\frac{p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)})}{1 - p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)})} \right) \in (-\infty, +\infty)$: log da chance de pertencimento a classe c :

Aplicando Regressão Linear para estimar o log das chances

$$\ln \left(\frac{p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)})}{1 - p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)})} \right) = \mathbf{X}^{(i)} \boldsymbol{\theta}$$

$$\frac{p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)})}{1 - p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)})} = e^{\mathbf{X}^{(i)} \boldsymbol{\theta}}$$

$$p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)}) = e^{\mathbf{X}^{(i)} \boldsymbol{\theta}} (1 - p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)}))$$

$$p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)}) + p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)}) e^{\mathbf{X}^{(i)} \boldsymbol{\theta}} = e^{\mathbf{X}^{(i)} \boldsymbol{\theta}}$$

$$p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)}) (1 + e^{\mathbf{X}^{(i)} \boldsymbol{\theta}}) = e^{\mathbf{X}^{(i)} \boldsymbol{\theta}}$$

$$p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)}) = \frac{e^{\mathbf{X}^{(i)} \boldsymbol{\theta}}}{(1 + e^{\mathbf{X}^{(i)} \boldsymbol{\theta}})} = \frac{1}{e^{-\mathbf{X}^{(i)} \boldsymbol{\theta}} (1 + e^{\mathbf{X}^{(i)} \boldsymbol{\theta}})}$$

$$p(\hat{\mathbf{y}}^{(i)} = c | \mathbf{X}^{(i)}) = \frac{1}{(1 + e^{-\mathbf{X}^{(i)} \boldsymbol{\theta}})}$$

Regressão Logística

- Ao adaptar a Regressão Linear para a tarefa de classificação binária obtemos a Regressão Logística
- $\hat{y} = \frac{1}{(1+e^{-\mathbf{x}^{(i)}\theta})}$
- Precisamos de uma nova função de custo
 - J_{SSR} não é convexo para a Regressão Logística
 - $J_{BCE} = -\frac{1}{N} \sum_{i=1}^N \left[\mathbf{y}^{(i)} \ln \hat{\mathbf{y}}^{(i)} + (1 - \mathbf{y}^{(i)}) \ln (1 - \hat{\mathbf{y}}^{(i)}) \right]$

Treinamento

Vamos derivar J_{BCE} em relação aos θ

$$\begin{aligned}\frac{\partial J_{\text{BCE}}}{\partial \theta} &= \frac{\partial (-y \ln \hat{y} - (1 - y) \ln (1 - \hat{y}))}{\partial \theta} \\&= -\frac{\partial (y \ln \hat{y})}{\partial \theta} - \frac{\partial ((1 - y) \ln (1 - \hat{y}))}{\partial \theta} \\&= -y \frac{\partial (\ln \hat{y})}{\partial \theta} - (1 - y) \frac{\partial \ln (1 - \hat{y})}{\partial \theta} \\&= -y \frac{\partial (\ln \hat{y})}{\partial \hat{y}} \frac{\partial \hat{y}}{\partial \theta} - (1 - y) \frac{\partial \ln (1 - \hat{y})}{\partial \hat{y}} \frac{\partial \hat{y}}{\partial \theta} \\&= -\mathbf{X}^\top \mathbf{y} \frac{1}{\hat{y}} \hat{y} (1 - \hat{y}) - \mathbf{X}^\top (1 - y) \left(-\frac{1}{1 - \hat{y}} \right) \hat{y} (1 - \hat{y}) \\&= \mathbf{X}^\top (-y + y\hat{y} + \hat{y} - \hat{y}y) \\&= \mathbf{X}^\top (\hat{y} - y)\end{aligned}$$

Regressão Logística

- Função de Custo J_{BCE} é convexa, mas não admite solução analítica
- O gradiente tem exatamente a mesma forma da Regressão Linear
- Gera uma fronteira de decisão linear (hiperplano) no espaço de atributos

Leituras Recomendadas

- Seções 3.1 a 3.3 de: Y.S. Abu-Mostafa, M. Magdon-Ismail e H.T. Lin.
Learning from Data: A Short Course. AMLBook.com, 2012. ISBN: 978-1-60049-006-4

Referências

- [1] Y.S. Abu-Mostafa, M. Magdon-Ismail e H.T. Lin. Learning from Data: A Short Course. AMLBook.com, 2012. ISBN: 978-1-60049-006-4.
- [2] U. Braga-Neto. Fundamentals of Pattern Recognition and Machine Learning. Springer International Publishing AG, 2020. ISBN: 978-3-030-27655-3.
- [3] T.M. Mitchell. Machine Learning. McGraw-Hill International Editions. McGraw-Hill, 1997. ISBN: 978-0-07-115467-3.