



Deep Learning
A-Z™
Python ile
Derin Öğrenme
Dr. Merve Ayıce KIZRAK

Deep Learning A-Z™

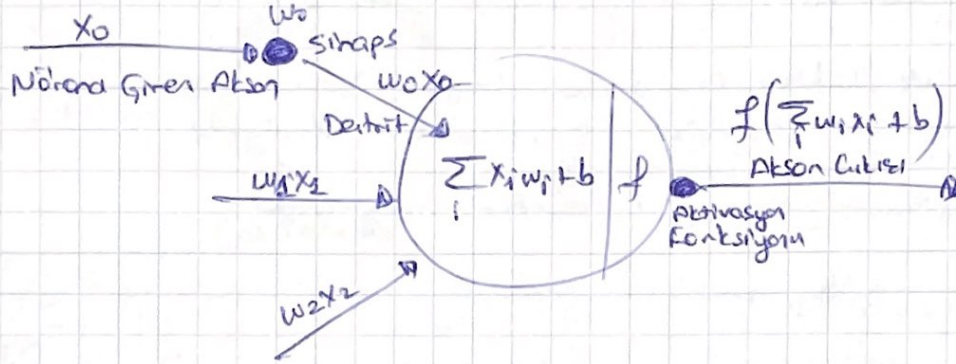
Python ile Derin Öğrenme

Dr. Merve Ayıce KIZRAK

UDENY

Doğukan Mehmet
KILINÇ

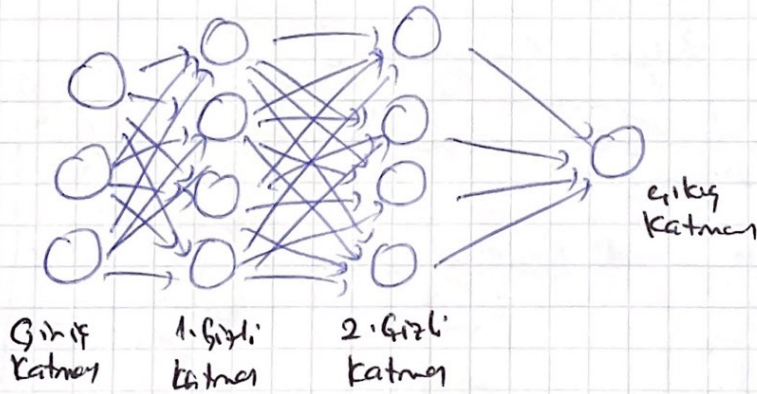
İnsanlardaki sinir ağının matematiksel modellenmesi aşağıdadır.



Buradaki fonksiyon çıkışı nihai çıkış veya bir başka fonksiyonun girişi olabilir.

$$\text{çıkış} \leftarrow \textcircled{y} = \textcircled{w} \times \textcircled{x} + \textcircled{b}$$

Ağırlıklar
girişler
bias



EĞİTİCİLİ (GÖZETİLMİ, SUPERVISED) ÖĞRENME

Sınıflandırma (Classification) Yöntemleri

- Destek Vektör Makineleri (Support Vector machines)
- Ayrıştırma Analizi (Discriminant Analysis)
- Naif Bayes (Naive Bayes)
- En Yakın Komşu (Nearest Neighbor)

Bağlılık (Regression) Yöntemleri

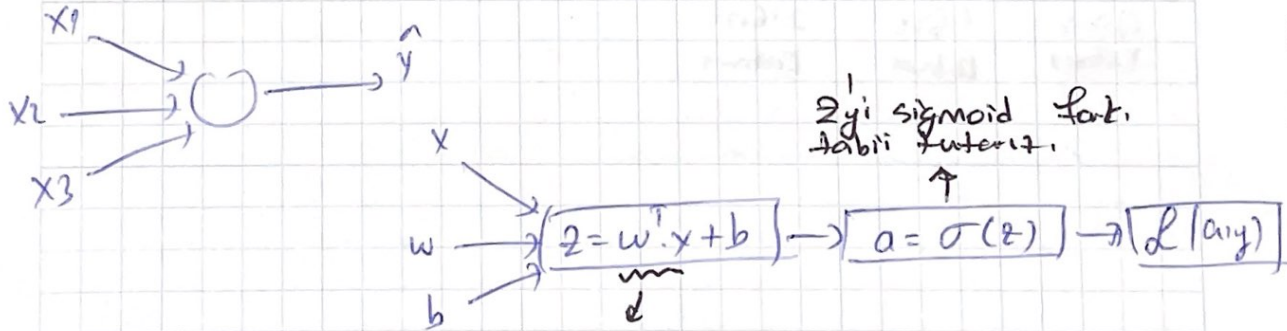
- Lineer Bağlılık
- Karar Ağaçları (Decision Trees)
- Yapay Sinir Ağları

EĞİTİCİSİZ (GÖZETİLSİZ, ~~UNSUPERVISED~~ REINFORCEMENT) ÖĞRENME

Kümeleme Yöntemleri

- K-Ortalama, Bulanık C-Ortalama
- Gauss Karışımı (Gaussian mixture)
- Saklı Markov Modeli (Hidden Markov model)
- Yapay Sinir Ağları

PEKİ ÖĞRENİM NASIL GERÇEKLEŞİYOR?

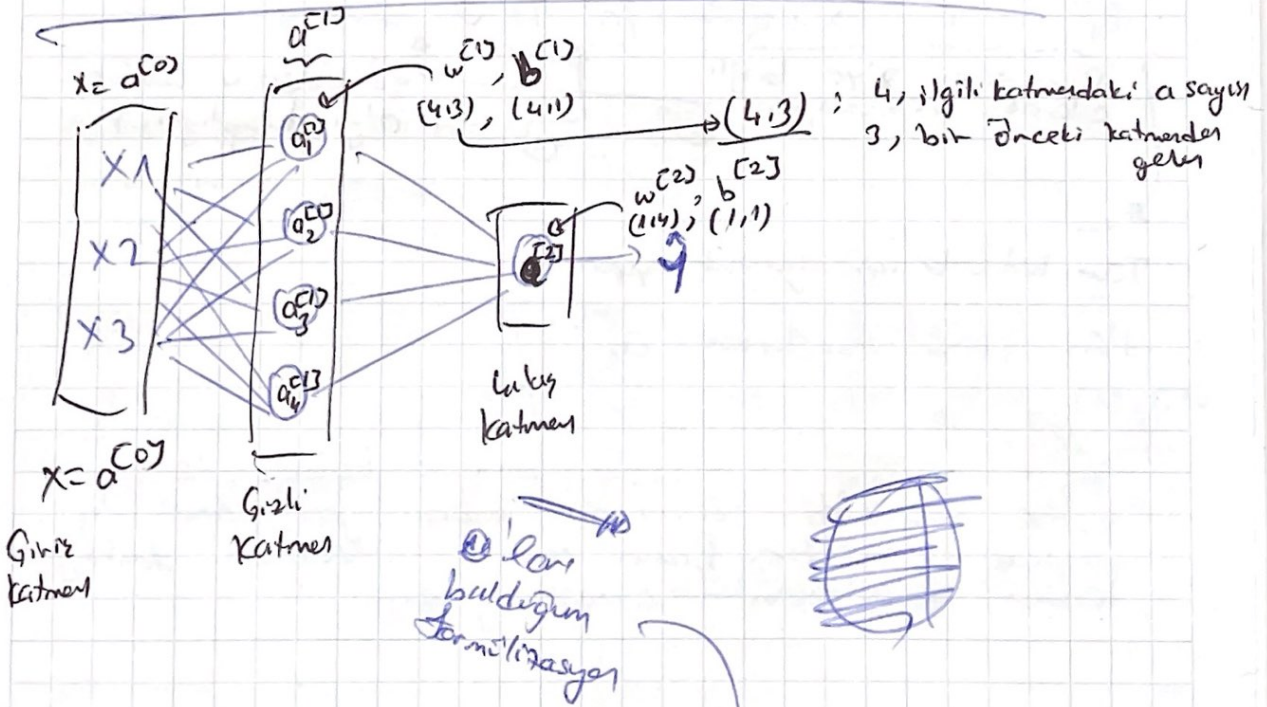
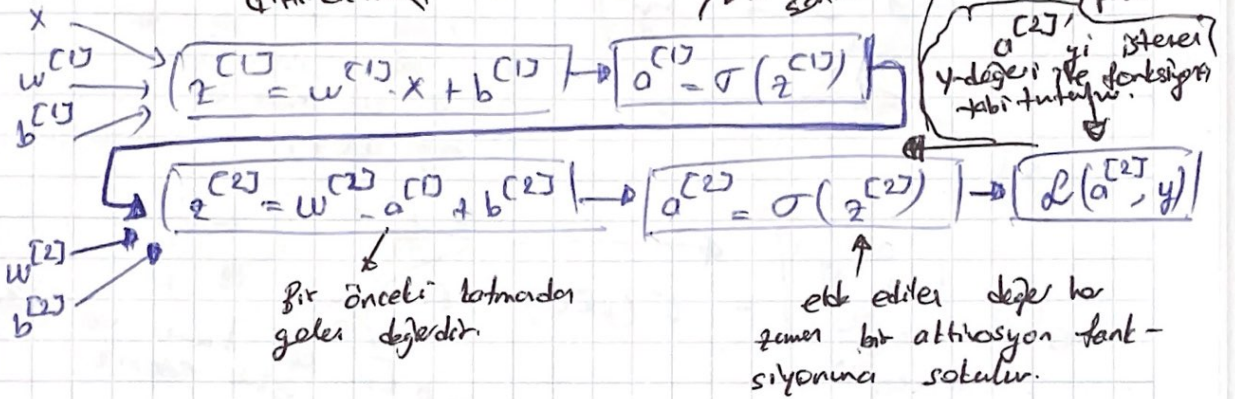
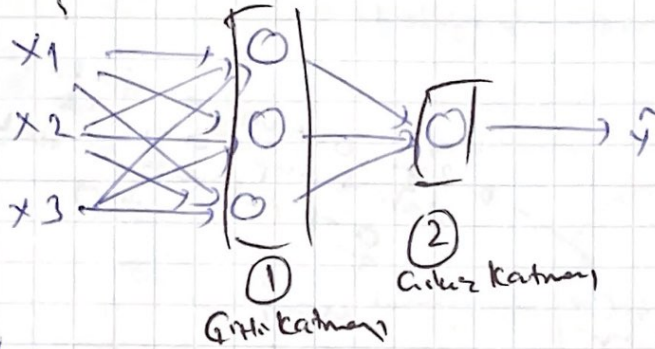


Görülele ağırlıklar hesaplanır.
Bunlar bir vektör olduğu için birinin
sütun sayısı ile diğerinin sütun sayısı
eşit olmalı. Bu yüzden ağırlığın
transpozunu alıyoruz.

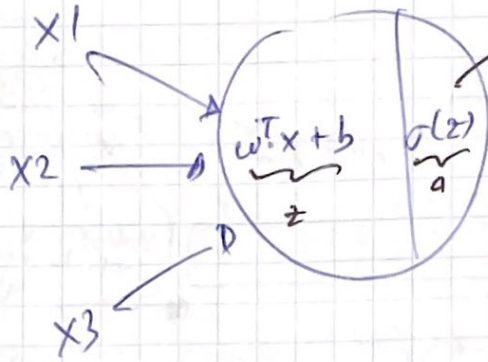
Girişlerdir. (Girişler zaman zaman sayılmaz)



...../...../20.....



a' 'leri bulduğum model ağırlıdır.



z 'yi aktivasyon fonksiyonuna
tabi tutup a 'yı elde ederiz.
 $a \rightarrow$ Aktivasyon fonk.
çıktı

$$z = w^T \cdot x + b$$

$$a = \sigma(z)$$

$$a^{[1]} = \begin{bmatrix} a_1^{[1]} \\ a_2^{[1]} \\ a_3^{[1]} \\ a_4^{[1]} \end{bmatrix} = \sigma(z^{[1]})$$

Budur!

En son bunu ve z_1, z_2, z_3 'ü
elde ederiz.

$$z_1^{[1]} = w_1^{[1]T} \cdot x + b_1^{[1]}, \quad a_1^{[1]} = \sigma(z_1^{[1]})$$

Bunadaki z 'ye bağlı
olacak aktivasyon

①

②

fonksiyonu sonucu
 a 'yı hesaplıyoruz

Tüm katmanlar için aynı işlemi yaparız.

ileri yönde beslerken

$$\begin{bmatrix} a_1^{[1]} \\ a_2^{[1]} \\ a_3^{[1]} \\ a_4^{[1]} \end{bmatrix}$$

buluruz. Fakat ileri
yönde besledikten sonra asıl modele geri yönde
yayılım yapmaktır. Bunun için de aktivasyon fonksiyon-
larının türevlenebilir olması gerekir.

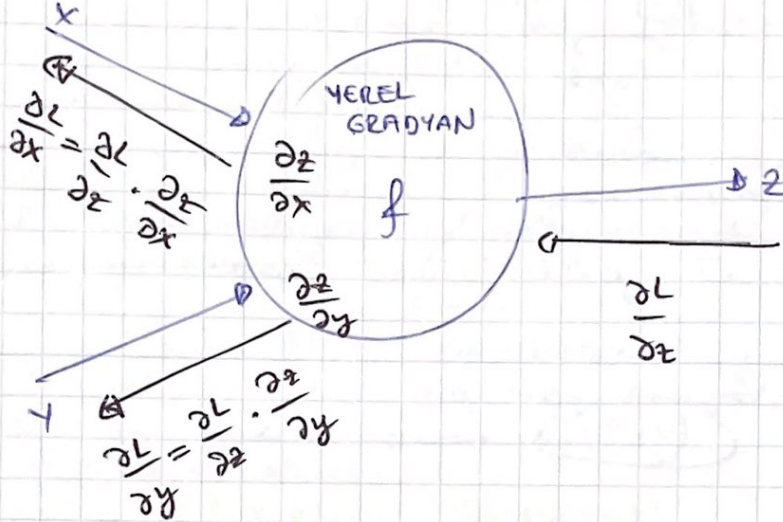
w: Ağırlıklar
x: Girdiler



...../...../20.....

GERİYE YAYILIM ALGORİTMASI

Giriş bilgilerine ulaşıldığında sistem sürekli ileri ve geri yönde beslemeye devam eder.



İleri geri yayılımı sona parametreler güncellenir. Fonksiyon minimize edilene kadar bu devam eder. Bu işlemin tamamına öğrenme işlemi denir.

İleriye Yayılım Algoritması (Feed Forward Propagation)

$$z^{[1]} = w^{[1]} \cdot x + b^{[1]}$$

$$A^{[1]} = g^{[1]}(z^{[1]})$$

$$z^{[2]} = w^{[2]} \cdot A^{[1]} + b^{[2]}$$

$$A^{[2]} = g^{[2]}(z^{[2]}) = \sigma(z^{[2]})$$

Geriyeye Yayılım Algoritması (Back Propagation)

$$dz^{[2]} = A^{[2]} - y$$

$$dw^{[2]} = \frac{1}{n} dz^{[2]} \cdot A^{[1]T}$$

$$dz^{[1]} = \underbrace{w^{[2]T} \cdot dz^{[2]}}_{(n^{[1]} \cdot m)} + \underbrace{g^{[1]'}(z^{[1]})}_{(n^{[1]} \cdot m)} \cdot z^{[1]}$$

$$dw^{[1]} = \frac{1}{n} dz^{[1]} \cdot x^T$$



FIGES

...../...../20.....

ARA HATIRLATMA

Bu şekilde hesaplanır

$$\begin{bmatrix} 2 & 3 & 4 \\ 2 & -7 & 4 \end{bmatrix} \times \begin{bmatrix} 3 & 4 & 5 \\ 1 & 4 & 4 \\ 2 & 1 & 4 \end{bmatrix} = \begin{bmatrix} (6+3+2) & (8+3+1) & (10+12+4) \\ (6-7+8) & (8-7+4) & (10-28+16) \end{bmatrix}$$

(2×3) (3×3)

Bu iki matrisi
çarpabilmek için
bu değerlerin eşit
olması gerekir.

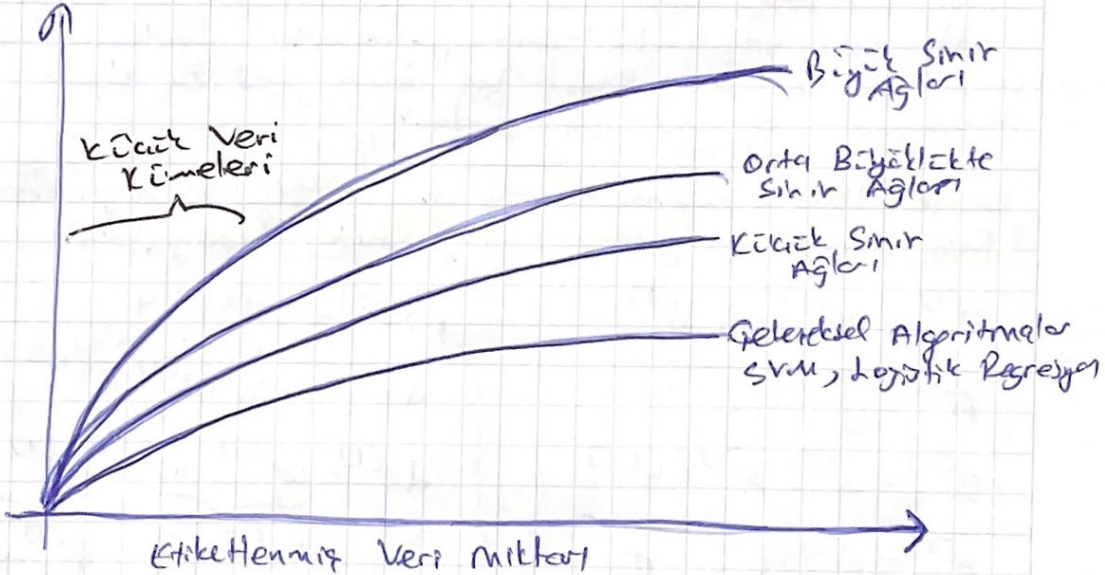
$$= \begin{bmatrix} 11 & 12 & 26 \\ 7 & 5 & -2 \end{bmatrix}_{2 \times 3}$$

Yeni elde edilen
matrisimiz bu olur.

ve yeni matrisimiz
bu boyuta eşit olur.

(2×3)

PERFORMANS



Veri Az ise Derin Sinir Ağı kullanmaya gerek yok çünkü işlem çok fazla. Gelişimsel Algoritmalar kullanılabilir. Ancak etiketlenmiş veri miktarı arttığında performans ciddi bir şekilde değiştiğinde Büyük Sinir Ağları tercih edilebilir.

HİPERPARAMETRE NEDİR?

Veriden öğrenen makine öğrenmesi ya da derin öğrenme modelleri tasarlamak ne olması gerektiği, belirli koşullar göz önüne alınarak modeli tasarlayan kişiye bırakılmış; probleme, veri setlerine ve donanımına da bağlı olarak değişiklik gösteren parametrelere **hiperparametre** denir. Tamamen değişken ve hiçbir zaman net bir şey söyleyemeyeceğimiz konudur.

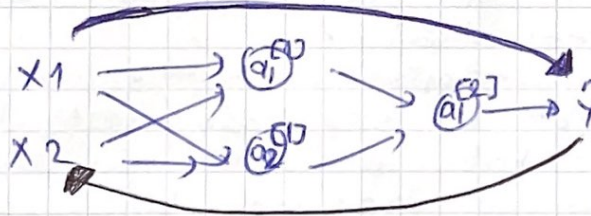
REGULARIZATION: Modele uygun parametrelerin hiperparametre optimizasyonu ile seçilmesidir.

NEDEN BASİT MODEL SEÇMELİYİZ?

- 1) Hesaplama karmaşıklığının düşük olması gerekiyor.
- 2) Eğitimin kolay olması, yer kaplamaması için.
- 3) Kolay açıklanabilir olması için.
- 4) Genelleştirilebilir olması.

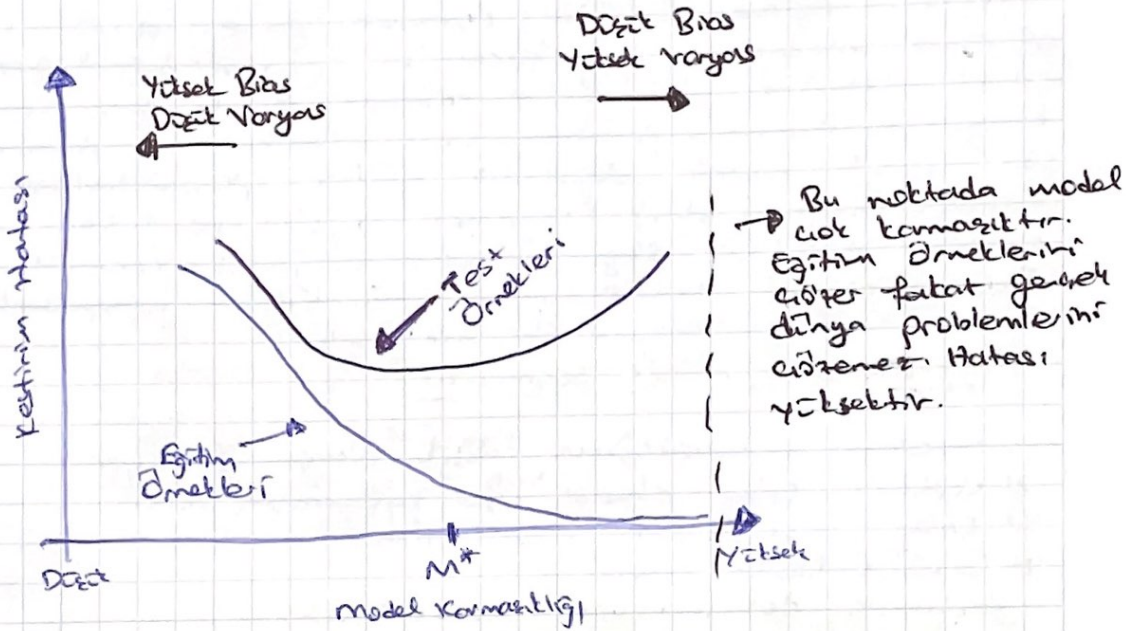
BASLANGIÇ AĞIRLIKLARININ BELİRLENMESİ

- Ağırlıklar sıfır (0) ile başlatılmamalı.
- Postgele küçük sayılarla başlatılmakta fayda var.



Her bir nöronda geriye yayılımında yeni değerle öğrenmesi için girişlerin sıfır değil küçük bir sayı ile başlatılmalıdır. Sıfırı eğer bir regresyon yapıyorsak, tek bir nöronumuz varsa sıfırla başlatmakta problem yoktur.

Bias - VARIANCE ilişkisi



Bias yüksek olduğunda, model çok azabilir ve genellebilir durumdadır. Fakat başarısı eğitim için de test için de çok yüksek değildir. Kestirim hatası yüksektir. Başarımı düşüktür.

Model karmaşıklığa bağladıkça hatalar azalır test ve eğitim birlikte. Bir süre sonra testin hatası ile eğitimin hatası birbirinden uzaklaşmaya başlar. Eğitimin hatası düşmeye devam eder, burada varyans yüksek olur çünkü test ve eğitim hatası arası mesafe azalır. Fakat bias düşük olur çünkü eğitimin hatası düşüktür.

Öyle bir model karmaşıklığı seçmeliyiz ki minimum olan bias - variance ilişkisi optimum noktada olsun.

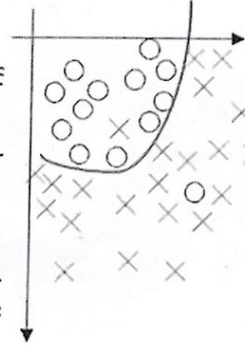
Model çok güzel
az karmaşık

yüksek bias

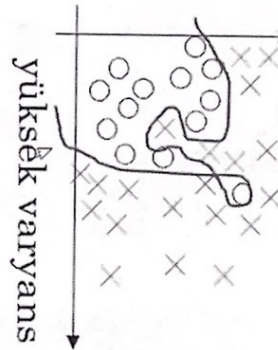
Az Uydurma
Under fitting

Bias = modelin tahminlerindeki sistematik hata.
Variance = eğitim setinde örnekleme gürültüsü tahminlerdeki gürültüye ne neden olur.

Bias vs. Variance

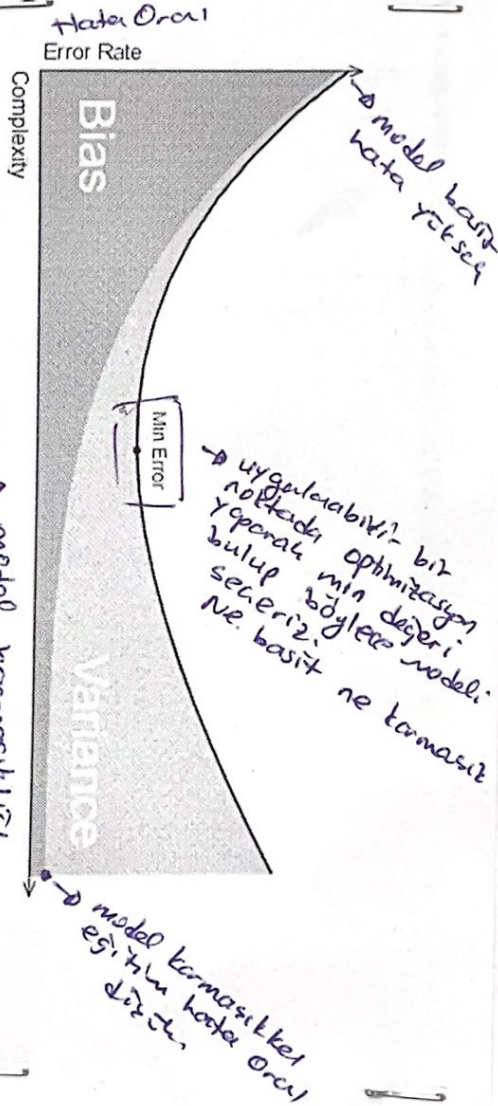


✓ "tam kararında"
Ne yüksek bias
Ne yüksek variance



yüksek varyans
Aşırı Uydurma
Over fitting

Model mükemmel seçilebilir.
Ama sadece bir
sınıflı veriye göre yapılmıştı



Model karmaşıklık artar

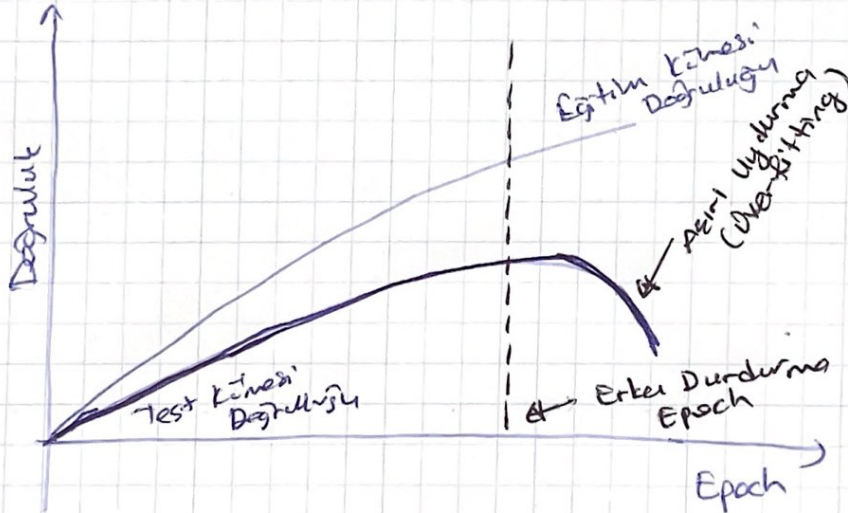
Düğüm SEYRELTME (DROPOUT)

Bu da regularizasyon yarı deneysel yöntemler-
ninde birisidir. Tanı bağlamı katmanlarda eşik değeri
altında kalan nöronların iptal edilmesi başarıyı
arttırabilir. Klem yükünü azaltır.

HEDEF: Modeli geliştirmeye yardımcı olur.

Genellikle eı sıra eklenen Dropout katmanında
0.50 kullanılır. (0.50 'sini' unit deriz)

Azı UYDURMA - AZ UYDURMA (Overfitting vs. UNDERFITTING)

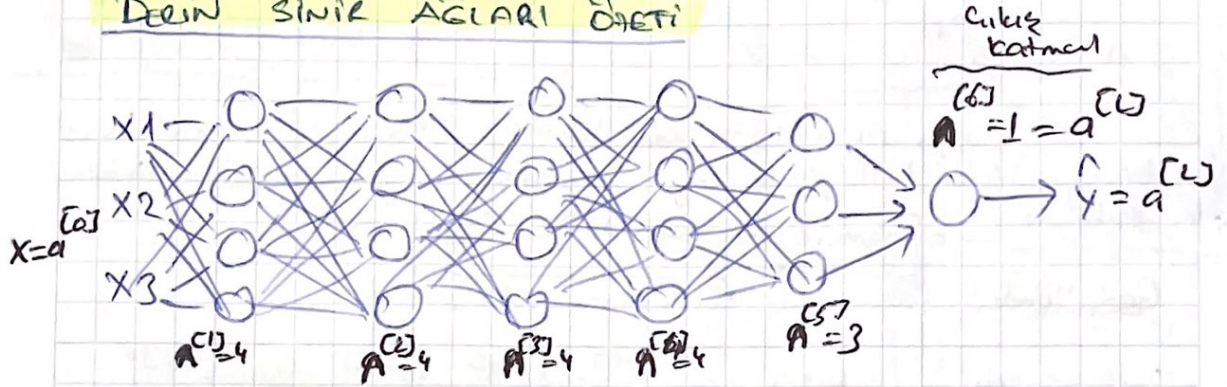


Aradaki mesafe artmaya başlayınca Early Stopping
(Erken Durdurma) yapılır ve devam eden epoch'lar
eakıştırılmamalıdır. Böylelikle bunun önüne geçilir
oluruz.

TRANSFER ÖĞRENME (TRANSFER LEARNING)

Özellikle çok katmanlı modellerde performansı ciddi derecede etkileyen bir düzenleme yöntemidir.

Derin SINIR AĞLARI ÖZETİ

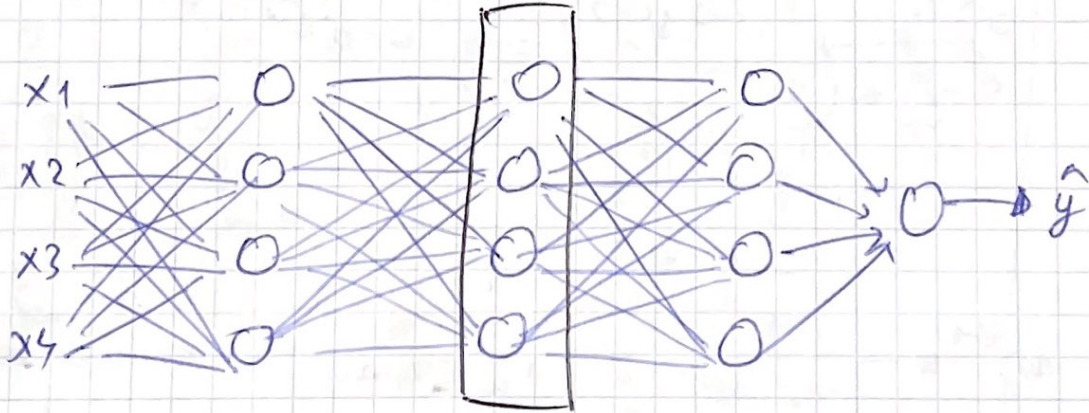


L : Katman Sayısı = 6

$n^{[l]}$: l . katmandaki nöron sayısı

$a^{[l]} = g^{[l]}(z^{[l]})$, $w^{[l]}$: l . katmandaki ağırlıklar.

Bu tip bir sınırlı ağırlık ileri ve geri yayılım için aşağıdaki yolu kullanırız



L . Katman için

L . katman için ileri yönde yayılım yaptığımızda bir önceki $(l-1)$ katmanın çıktıları bizim girişimiz olacaktır.

İleri Yönde: → Bir önceki katmandan gelen bilgiler.

Giriş: $a^{[l-1]}$, Çıkış: $a^{[l]}$, ARA: $z^{[l]}$

$$z^{[l]} = w^{[l]} \cdot a^{[l-1]} + b^{[l]} \quad \left(\begin{array}{l} (x \text{ girişleri}) \cdot (\text{kollarındaki ağırlıklar}) \\ + (\text{bias değerleri}) \end{array} \right)$$

$$a^{[l]} = g(z^{[l]})$$

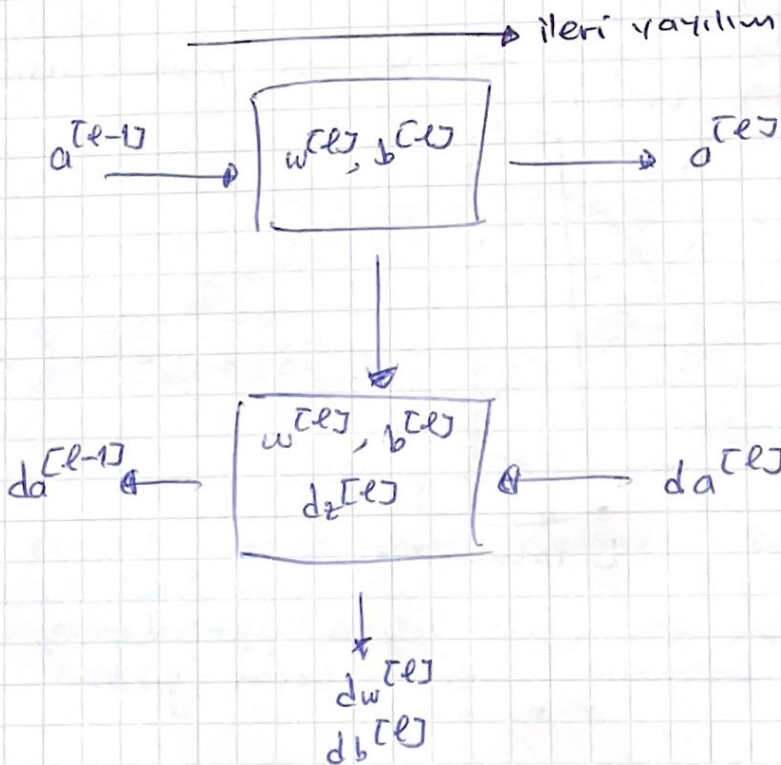
g : Aktivasyon fonksiyonuna tabi tuttuğumuzda aynı kopyasını. Bir önceki katmandan gelen girişleri ifade eder.

Sigmoid, Hiperbolik Tangent, Relu da bilinir.

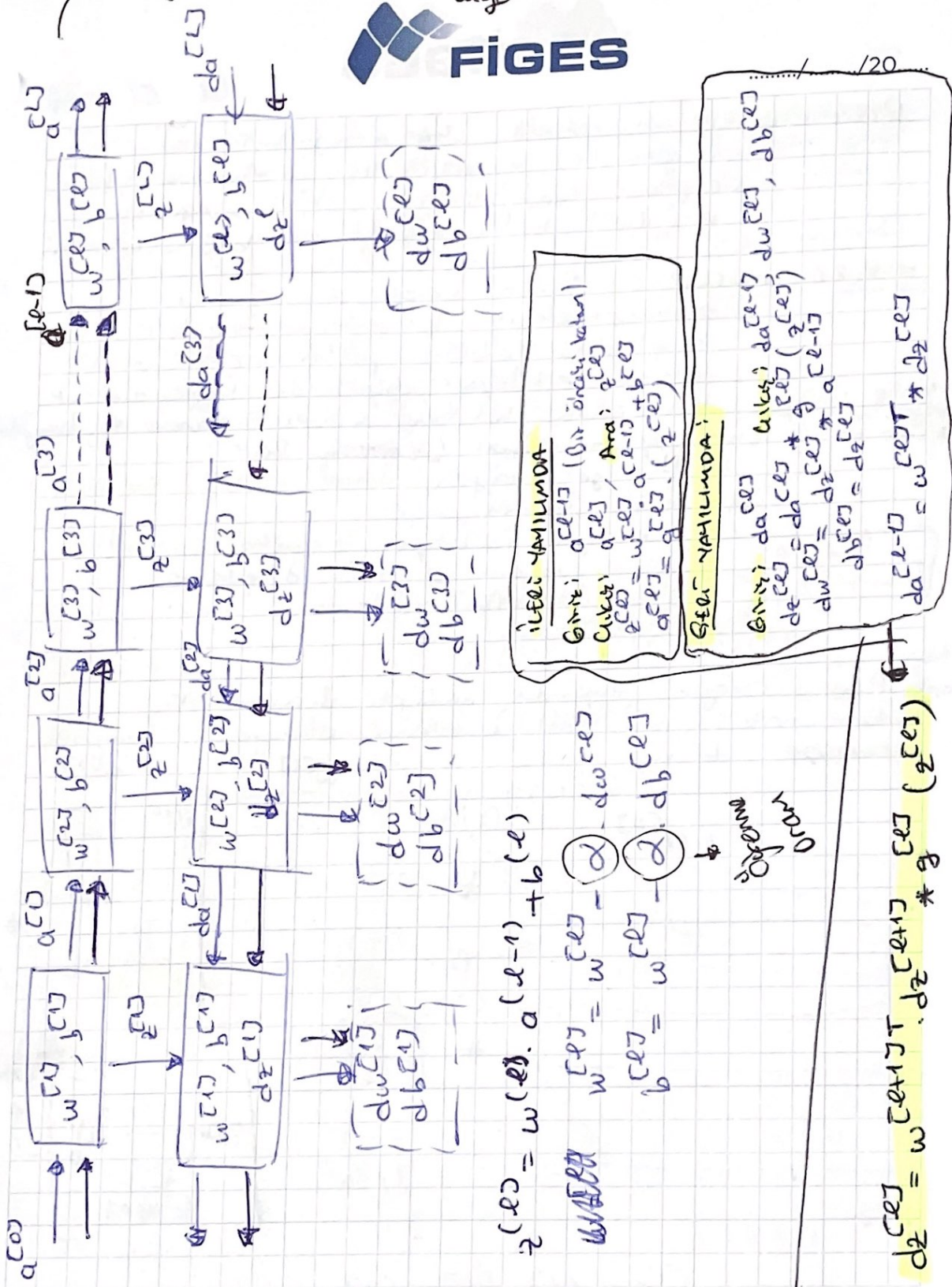
Geri Yönde: Bu safher türevlerini alınız.

Giriş: $da^{[l]}$, Çıkış: $da^{[l-1]}$, $dw^{[l]}$, $db^{[l]}$, ARA: $z^{[l]}$

L. Katman için:



Elde edilen Loss Fonksiyonunu minimize etmek için geriye yayılım algoritmasını kullanırız.



(*) → Çarpma değil
evrilim referidir.

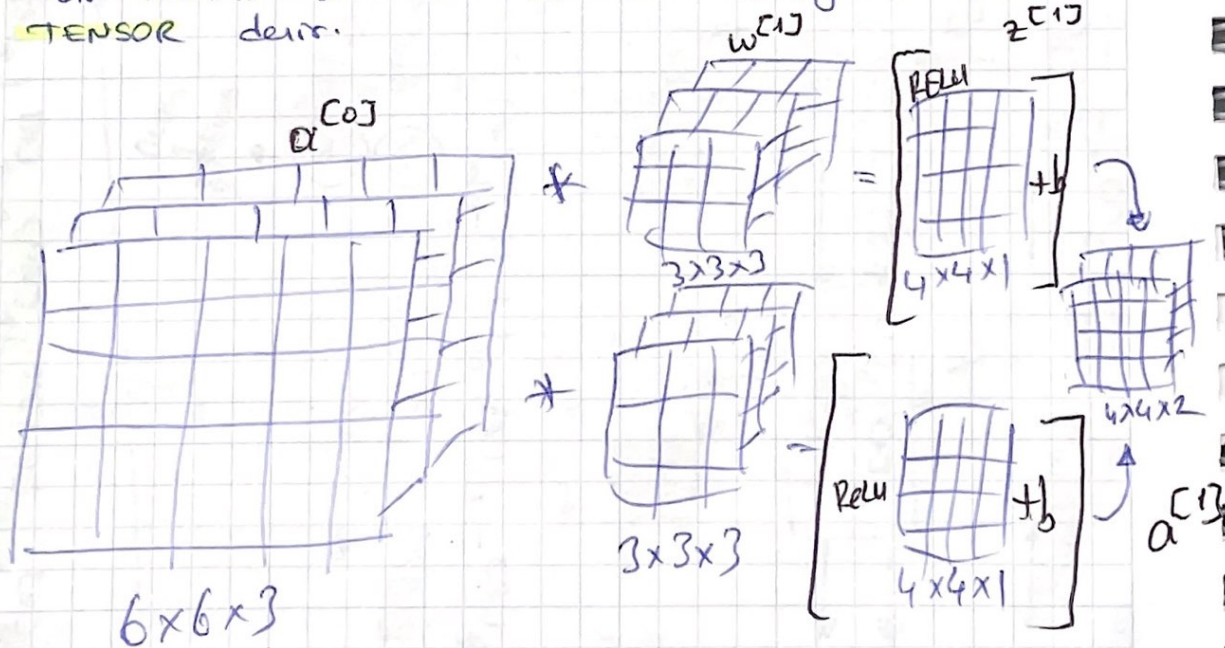


04.01.2024

PARAMETRELER: Modellerde her zaman bulunması
gerekir, hesaplanması gereken ağırlıklar
ve bias değerleridir.
 $\rho (W^{[1]}, b^{[1]}, W^{[2]}, b^{[2]})$ W : Ağırlıklar
 b : Bias Değerleri

HİPERPARAMETRELER: Bizim karar vereceğimiz ve
parametrelerin hesaplanmasına katkı
sağlayan, modelin performansını direkt
olarak etkileyen değerlerdir. Değişkendir ve
her modelde her zaman farklı performans sağlanabilir.
• Öğrenme Oranı (Learning Rate)
• İterasyon Sayısı
• Girdi Katman Sayısı
• Aktivasyon fonksiyonu seçimi
• Kırma Boyutu ve diğer düzenleme
(regularization) yöntemleri

Anaak iterasyon sayısında maliyeti düşürmektir.
Bir matris bir çok kısımdan oluşuyorsa buna
TENSOR denir.



$$z^{[1]} = W^{[1]} \cdot a^{[0]} + b^{[1]}$$
$$a^{[1]} = g(z^{[1]})$$

Giriş boyutu büyüdükçe filtrelerin ne kadar kullanıldığına bağlı olarak parametreler değişir. Giriş boyutuna göre değil. Ancak sadece filtrelerin girişe uygularken daha fazla işlem yapmasını gerektiriyor.

Bu filtreler birer birer örneklilikleri hesaplanmaktadır.

Filtre Boyutu: $f^{[l]}$
 Pixel Eklendirme: $p^{[l]} \rightarrow \text{Padding}$
 Adım Kaydırma: $s^{[l]} \rightarrow \text{stride}$
 Filtre (kanal) Sayısı: $n_c^{[l]} \rightarrow \text{N-channel}$

Her bir filtrenin boyutu: $f^{[l]} \times f^{[l]} \times n_c^{[l-1]}$

Aktivasyon fonksiyonu: $a^{[l]} \rightarrow n^{[l]} \times n_w^{[l]} \times n_c^{[l]}$

Ağırlıklar

Bias

Giriş: $n_H^{[l-1]} \times n_W^{[l-1]} \times n_c^{[l-1]}$

Çıkış: $n_H^{[l]} \times n_W^{[l]} \times n_c^{[l]}$

$$n_{H,W}^{[l]} = \left\lceil \frac{n_H^{[l-1]} + 2p - f^{[l]}}{s^{[l]}} + 1 \right\rceil$$

Giriş görüntüsünün kanal sayısı ile aynı oldu

kanal (Channel)

çıkış boyutu ile aynı olmadı

Bir Evrimsel Sınırlı Ağ ve Gerekli Katmanlar

- Evrimsel Katman (Aktivasyon Fonksiyonu, Bias Değeri)
- Ortalama Katman (maksimum ya da Ortalama Ortalama)
- Tem / Tem Bağlantı Katman (Klasik Yöney Sınırlı Ağ Bağlantıları)

ÖRNEK

1. Adım
Girte Gözetimi
Bağlantı

$$a^{[0]} \Rightarrow 27 \times 27 \times 3 \text{ olsun}$$

$$L \Rightarrow n_H^{[0]} = n_W^{[0]} = 27$$

$$n_C^{[0]} = 3$$

$$f^{[0]} = 3$$

$$p^{[0]} = 0$$

$$s^{[0]} = 1$$

olan bir
filtre uyguladığında
ve 10 filtre uygulanca

Yeni
Girdi

$$\frac{n + 2 \cdot p - f^2}{s} + 1$$

$$= \frac{27}{1} + 1 = 28$$

2. Adım

$a^{[1]}$

$25 \times 25 \times 10$ matrisine

$$\begin{cases} n_H^{[1]} = n_W^{[1]} = 25 \\ n_C^{[1]} = 10 \end{cases}$$

$$f^{[1]} = 5$$

$$p^{[1]} = 0$$

$$s^{[1]} = 2$$

değerli filtre uygulayalım
20 tane filtre uygulayalım.

$$\left\lceil \frac{n + 2p - f}{s} + 1 \right\rceil \text{ de}$$

$$\frac{25 + 0 - 5}{2} + 1 = 11$$

$(25 \times 25 \times 10)$ f -10p
yeni
matris elde ederiz.

$a^{[2]}$ 20 Filtre
uyguladık

$$11 \times 11 \times 20$$

matris halinde
ikinci katman
da bulunur.

3. Adım: $11 \times 11 \times 20$ 'lık boyuttaki matrise

$$n_H = n_W = 11$$

$$n_C = 20$$

$$f^{[3]} = 5$$

$$p^{[3]} = 0$$

$$s^{[3]} = 2$$

ve 40 filtre uygulayalım.

$$\left\lfloor \frac{n + 2p - f}{s} \right\rfloor + 1 \rightarrow \frac{11 + 2 \cdot 0 - 5}{2} = 3 + 1 = 4$$

Yeni matris boyutu $\rightarrow 4 \times 4 \times 40$

4. Adım: $4 \times 4 \times 40$, benim son katmanım ise artık bunu vektör haline getirmişim ve $4 \times 4 \times 40 = 640$ satırlık bir vektör olur. Bu da çıkış nöronuna ulaştırılır. Burası da benim elde ettiğim kestirim değeri olur.

Ortaklaşma katmanında kanal sayısı değişmez. Sadece yükseklik ve genişlik bakımından boyut azaltma gerçekleştirilir.

Yapılan işlemler bu adımla sonrasında yükseklik ve genişlik azalırken kanal sayısı artacaktır.

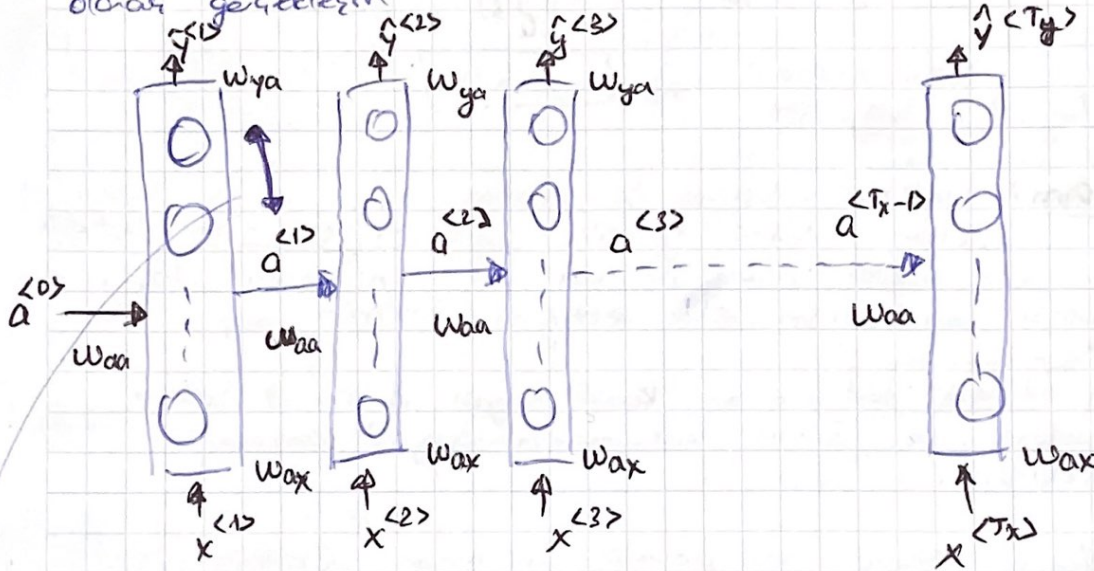
$$n_H \downarrow \quad n_W \downarrow \quad n_C \uparrow$$

ÖĞRETİLEBİLİR SİNİR AĞLARI (Recurrent Neural Networks - RNN)

Zaman serisi şeklinde ifade edilen problemleri
çözüm diyebiliriz.

Özellikle ses tanıma, ses ile ilgili problemler.
Çünkü ses zaman boyunca süregelen bir veridir
ve bu veri öğrenilebilir sinir ağı ile
çözülebilir bir veridir.

RNN'de kısımlar her zaman bir önceki kısımla bağlı
olarak geçerlidir.



$$a^{(0)} = 0$$

bir önceki katman

$$a^{(1)} = g_1 (W_{aa} \cdot a^{(0)} + W_{ax} \cdot x^{(1)} + b_a)$$

girişteki değerler

a için bias

Bu bir aktivasyon fonksiyonudur.
(ReLU, RNN'de pek tercih edilmez.)
(Genellikle hiperbolik tangant tercih edilir.)

$$y^{(1)} = g_2 (W_{ya} \cdot a^{(1)} + b_y)$$

Aynı aktivasyon fonksiyonu
olmak zorunda değildir. Ancak
aynı da olabilir.

NOTASYON

(W_{ya}) y: Hesapladığımız
a: Alınan
seçilmiş
yazılır

önce hesaplanır
sonra cevaplar yazılır.

Genelleme Yapacak Olursak:

$$\begin{aligned} a^{(t)} &= g(w_{aa} \cdot a^{(t-1)} + w_{ax} \cdot x^{(t)} + b_a) \\ \hat{y}^{(t)} &= g(w_{ya} \cdot a^{(t)} + b_y) \end{aligned} \quad \left\{ \begin{array}{l} t: \text{Ait olduğumuz an.} \\ t-1: \text{Bir önceki an.} \end{array} \right.$$

Genelleme a 'lar hesaplanırken $\Rightarrow$ Hiperbolik Tangent veya Relu tercih edilir. Ancak Relu çok nadir tercih edilir.

Genelleme $\hat{y}$ (çıktılar) hesaplanırken $\Rightarrow$ Sigmoid tercih edilir.

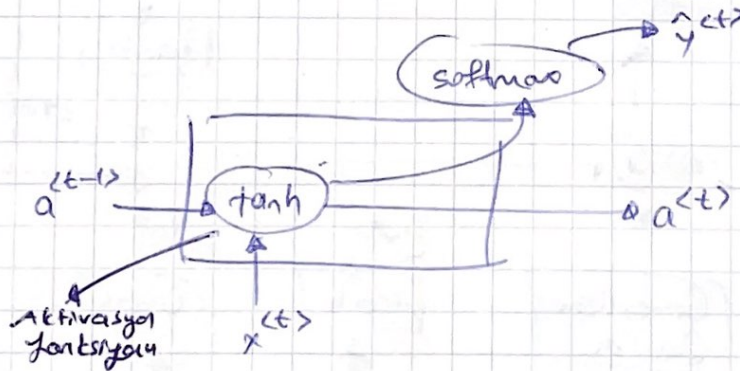
w_{aa}, w_{ax} 'ler $\rightarrow$ MATRİS
 a ve x 'ler $\rightarrow$ VEKTÖR'dür.

Notasyonu daha da sistematikleştirelim:

$$\begin{aligned} a^{(t)} &= g(w_a [a^{(t-1)}, x^{(t)}] + b_a) \quad \text{ve} \\ y^{(t)} &= g(w_y [a^{(t)}], b_y) \quad \text{şeklinde de ifade edebiliriz.} \end{aligned}$$

! Yine de ileri yayılım!

GRU (Gated Recurrent Units) (Geciktirmiş Öğrenilebilir Birimler.)



BELLEK

$$c^{(t)} = a^{(t)}$$

$$a^{(t)} = g(w_a [a^{(t-1)}, x^{(t)}] + b_a)$$

Γ : Unutma Değeri

LSTM (Long Short Term Memory) (Uzun Kısa Vadeli Bellek)

GPU'nun Denklemleri

İlgili

$$\tilde{c}^{(t)} = \tanh(w_c[\Gamma^{(t-1)} * c^{(t-1)}, x^{(t)}] + b_c) \rightarrow \tilde{c}^{(t)} = \tanh(w_c[a^{(t-1)}, x^{(t)}] + b_c)$$

Güncelleme Geçidi

$$\Gamma_u = \sigma(w_u[c^{(t-1)}, x^{(t)}] + b_u) \rightarrow \Gamma_u = \sigma(w_u[a^{(t-1)}, x^{(t)}] + b_u)$$

İlgililik Geçidi

$$\Gamma_f = \sigma(w_f[c^{(t-1)}, x^{(t)}] + b_f) \rightarrow \text{Unutma Geçidi}$$

$$\Gamma_f = \sigma(w_f[a^{(t-1)}, x^{(t)}] + b_f)$$

Çıkış

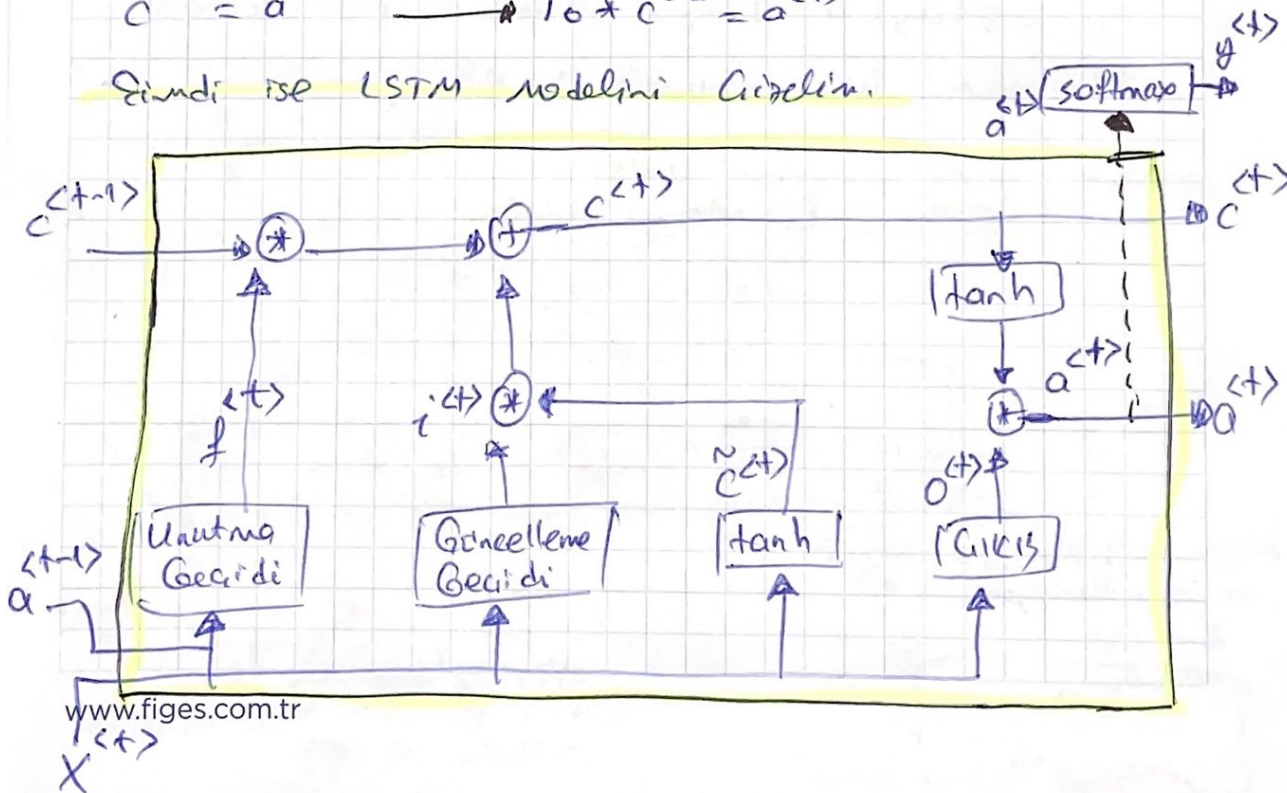
$$\Gamma_o = \sigma(w_o[a^{(t-1)}, x^{(t)}] + b_o)$$

$$c^{(t)} = \Gamma_u * \tilde{c}^{(t)} + (1 + \Gamma_u) * c^{(t-1)}$$

$$\hookrightarrow c^{(t)} = \Gamma_u * \tilde{c}^{(t)} + (\Gamma_f) * c^{(t-1)}$$

$$c^{(t)} = a^{(t)} \quad \rightarrow \quad \Gamma_o * c^{(t)} = a^{(t)}$$

Şimdi ise LSTM modelini çizelim.



Bu haritayı formülleştirirsek:

$$\tilde{c}^{(t)} = \tanh(W_c [a^{(t-1)}, x^{(t)}] + b_c)$$

$$\Gamma_u = \sigma(W_u [\underline{a}^{(t-1)}, x^{(t)}] + b_u) \leftarrow \text{Güncelleme geçidi olduğu yer}$$

$$\Gamma_f = \sigma(W_f [\underline{a}^{(t-1)}, x^{(t)}] + b_f) \leftarrow \text{Unutma geçidi olduğu yer}$$

$$\Gamma_o = \sigma(W_o [\underline{a}^{(t-1)}, x^{(t)}] + b_o) \leftarrow \text{Çıkış geçidi olduğu yer}$$

$$c^{(t)} = \frac{\Gamma_u * \tilde{c}^{(t)}}{\Gamma} + \frac{(\Gamma_f) * c^{(t-1)}}{\Gamma}$$

Güncelleme değeri ve e'ye aday olan + c'ndeki değer alınıyor.

Unutma geçidinden çıkan değeri de bir önceki t c'ndeki çıkan c değeri alınıyor.

$$\Gamma_o * c^{(t)} = a^{(t)}$$

Bunlar birbirine ard arda eklenebilir. Birinin çıkışı diğerinin girdisi olur.