

# پیشنهاد توابع فعال‌ساز بازه‌ای در شبکه عصبی بر پایه توابع شعاعی برای پیش‌بینی سیستم‌های غیر خطی پویا

الله یار ظهوری زنگنه<sup>۱</sup>، محمد تشنه‌لب<sup>۲</sup>، مجتبی احمدیه خانه‌سر<sup>۳</sup>

<sup>۱</sup> دانشجوی دکترای مهندسی کامپیوتر- هوش مصنوعی، گروه کامپیوتر، دانشگاه آزاد اسلامی واحد علوم و تحقیقات تهران،

z.zangeneh@gmail.com

<sup>۲</sup> استاد، قطب علمی کنترل صنعتی، دانشکده مهندسی برق، گروه کنترل، دانشگاه صنعتی خواجه نصیرالدین طوسی، teshnehlab@eetd.kntu.ac.ir

<sup>۳</sup> استادیار، دانشکده برق و کامپیوتر، گروه مهندسی قدرت و کنترل، دانشگاه سمنان، ahmadih@semnan.ac.ir

تاریخ دریافت مقاله ۱۳۹۴/۷/۲۳، تاریخ پذیرش مقاله ۱۳۹۴/۱۰/۴

**چکیده:** «شبکه عصبی بر پایه توابع شعاعی» یک تقریب گر عمومی می‌باشد. در این مقاله «تابع فعال‌ساز گرانولی» برای بهبود یادگیری این شبکه در نویزی پیشنهاد می‌گردد که یک تابع گاوسی با «انحراف استاندارد بازه‌ای و میانگین ثابت» است و به آن «تابع فعال‌ساز بازه‌ای» نیز گفته می‌شود. در لایه میانی این شبکه، سه پارامتر وابسته به توابع فعال‌ساز گرانولی آموزش می‌بینند که «مرکز توابع فعال‌ساز گرانولی» که مرکز دسته نامیده می‌شود، کران پائین انحراف استاندارد و کران بالای انحراف استاندارد این توابع می‌باشند. در لایه خروجی دو پارامتر دیگر یعنی «مرکز وزن‌های بازه‌ای» و «بازه این وزن‌ها» آموزش می‌بینند. برای آموزش این پارامترها از روش «الگوریتم خوشه‌بندی K-Means» استفاده شده است. در این روش، آموزش شبکه در راستای «گرانوله‌سازی پائین به بالا» می‌باشد که در آن بردارهای ورودی به شکل گرانول‌های بزرگتر در لایه میانی خوشه‌بندی می‌گردند. از روش «گرادیان نزولی» نیز برای آموزش پارامترهای شبکه استفاده شده و نتایج با روش جدید مقایسه گردیده است. عملکرد این شبکه با شناسایی «یک سیستم غیر خطی پویای U شکل با پنج ورودی» و «پیش‌بینی سری زمانی آشوب مکی گلاس» در شرایط نویزی و بدون نویز سنجیده می‌شود. از نتایج معلوم می‌گردد که استفاده از تابع فعال‌ساز گرانولی در ساختار شبکه عصبی RBF باعث کاهش حساسیت به تغییرات ورودی شده و عملکرد آن در شرایط نویزی بهبود می‌یابد.

**کلمات کلیدی:** شبکه عصبی بر پایه توابع شعاعی گرانولی<sup>۱</sup>، تابع فعال‌ساز گرانولی<sup>۲</sup>، داده‌های نویزی، انحراف استاندارد بازه‌ای، سیستم‌های غیر خطی پویا و توابع آشوب.

## Proposing Interval Activation Functions in Radial Basis Function Neural Network to Predict Nonlinear Dynamic Systems

Allahyar Zohoori Zangeneh, Mohammad Teshnehlab, Mojtaba Ahmadih Khanesar

**Abstract:** A Radial Basis Function Neural Network (RBFNN) is a general approximator. In this paper a granular activation function is proposed to improve its learning under the noisy conditions. The granular activation function is also named the interval activation function and it is typically the Gaussian function which benefits from having a fixed mean and an uncertain standard deviation. The hidden layer of the proposed network has a total of three parameters to train that it consists the means, the lower bounds of the standard deviations and the higher bounds of the standard deviations of the Gaussian functions. The output layer parameters for training are the means of the interval weights and the intervals of the weights. "K-Means clustering algorithm" method is used to train

<sup>1</sup> Radial Basis Function Neural Network (RBFNN)

<sup>2</sup> Granular activation functions

<sup>3</sup> Bottom-up granulation

<sup>4</sup> Nonlinear dynamic system with multiple time delays

<sup>5</sup> Mackey glass chaotic time series

<sup>6</sup> Granular Radial Basis Function Neural Network (GRBFNN)

these parameters. The purpose of the above learning method is regarded as one of the granular method presenting the bottom-up granulation which causes the input vectors clustered in the larger granules in the hidden layer. Gradient descend method is also used to train these parameters to compare with this novel method. The structure is tested with or without noisy data to identify a nonlinear dynamic system with multiple time delays and to predict a chaotic model, Mackey-Glass. It has been shown that the sensibility related to input alterations reduces because of using the granular activation function in RBF Neural Network structure and the response of Granular RBF Neural Network with noisy data is better than RBF Neural Network.

**Keywords:** Granular Radial Basis Function Neural Network (GRBFNN), Granular activation function, Noisy data, Interval standard deviation, Nonlinear dynamic systems and Chaotic Models.

## ۱- مقدمه

یک تابع شعاعی<sup>۱</sup> تابعی است که مقدار آن فقط به فاصله از یک نقطه که مرکز دسته می‌باشد وابسته است. این توابع در تخمین<sup>۲</sup> [۲]، پیش‌بینی سری‌های زمانی<sup>۳</sup> [۱، ۴۳] و کنترل<sup>۴</sup> [۲۲] مورد استفاده قرار می‌گیرند. در یک شبکه عصبی، از این توابع می‌توان؛ به عنوان تابع فعال‌سازی نرون استفاده کرد. یک شبکه عصبی بر پایه توابع شعاعی در حالت کلی شامل سه لایه می‌باشد، لایه ورودی، لایه میانی یا همان لایه پنهان<sup>۵</sup> و لایه خروجی<sup>۶</sup> [۱، ۲۲-۲۶].

در این مقاله یک شبکه عصبی RBF گرانولی، برای پردازش سیستم‌های همراه با نویز معرفی می‌گردد. توابع دارای «میانگین بازه‌ای و انحراف استاندارد ثابت»<sup>۷</sup> و یا «انحراف استاندارد بازه‌ای و میانگین ثابت»<sup>۸</sup>، دو نوع تابع فعال‌ساز گرانولی می‌باشند [۶-۷، ۲۳]. یک شبکه عصبی RBF گرانولی، توابع فعال‌ساز گاوسی گرانولی با «انحراف استاندارد بازه‌ای و میانگین ثابت» را، در یک سیستم عصبی وارد می‌نماید تا مقاومت در برابر نویز آن را افزایش دهد [۶-۷، ۱۰-۱۱، ۱۸، ۲۰].

یک سیستم عصبی گرانولی روشی مناسب‌تر نسبت به دیگر شبکه‌های عصبی در بررسی عدم قطعیت‌ها می‌باشد [۵]. زیرا در صورت وجود نویز در مقادیر ورودی، پایداری سیستم را در مقابل آن افزایش می‌دهد [۶-۷، ۱۰].

یک شبکه عصبی RBF گرانولی، می‌تواند به عنوان کلاسی از شبکه‌های تطبیق‌پذیر محسوب گردد. شبکه‌های عصبی تطبیق‌پذیر، توسعه یافته شبکه‌های پیش‌رو هستند که تابع فعال‌ساز نرون‌های آن می‌تواند به هر فرمی باشد یعنی لزومی ندارد به صورت سیگموید، سینوسی یا گاوسی باشد [۲]. این کلاس از شبکه‌های تطبیق‌پذیر در کاربردهای شناسایی،

پیش‌بینی و کنترل سیستم‌ها استفاده شده و قابلیت خود را به خوبی نشان داده است. به سبب قابلیت‌های خوب این شبکه‌ها، می‌توان از آن به عنوان یک مدل کننده خطی سازی محلی و یک روش خطی سازی متعارف در تخمین حالت و کنترل نیز نام برد [۲].

روش آموزش پارامترها در شبکه عصبی RBF گرانولی یک مسأله با اهمیت است. یکی از این روش‌ها بر مبنای «گرادیان نزولی»<sup>۷</sup> بنا شده است، که محاسبه مشتق در بعضی مراحل آن در قانون زنجیره‌ای<sup>۸</sup> بسیار دشوار است. از طرفی به کار بردن آن ممکن است ما را در مینیمم محلی<sup>۹</sup> گرفتار سازد [۳۵، ۳۶]. پیچیدگی<sup>۱۰</sup> الگوریتم گرادیان نزولی برای آموزش پارامترها در شبکه عصبی RBF و شبکه عصبی RBF گرانولی فارغ از تعداد ورودی‌ها از درجه ۴ یعنی  $O(m \times n \times T \times Epoch)$  می‌باشد؛ که  $m$  تعداد نرون‌های لایه میانی،  $n$  ابعاد ورودی‌ها،  $T$  تعداد نمونه‌های ورودی و  $Epoch$  تعداد دفعات تکرار الگوریتم است [۳۵]. در این پژوهش سعی شده است یک روش دیگر که بتواند پارامترها را سریع‌تر و آسان‌تر از «گرادیان نزولی» آموزش دهد معرفی گردد. زیرا در روش «گرادیان نزولی»، همگرایی پارامترها به شدت وابسته به نرخ آموزش<sup>۱۱</sup> مناسب است که اغلب یافتن آن بسیار دشوار می‌باشد [۳۶].

در روش پیشنهادی این مقاله که «الگوریتم خوشه‌بندی K-Means»<sup>۱۲</sup> نامیده می‌شود همگرایی پارامترها سریع‌تر و در تعداد دفعات تکرار (مرحله)<sup>۱۲</sup> کمتری [۲۳، ۳۷، ۴۴] صورت می‌گیرد. پیچیدگی<sup>۱۰</sup> الگوریتم خوشه‌بندی K-Means برای آموزش پارامترها در شبکه عصبی RBF گرانولی بدون توجه به تعداد ورودی‌ها از درجه ۳ یعنی  $O(m \times T \times Epoch)$  می‌باشد؛ که  $m$  تعداد خوشه‌ها (یا همان تعداد نرون‌های لایه میانی)،  $T$  تعداد کل اشیاء (یا همان تعداد نمونه‌های ورودی) و  $Epoch$  تعداد دفعات تکرار الگوریتم است [۳۷]. آموزش پارامترهای لایه میانی بوسیله روش «الگوریتم خوشه‌بندی K-Means»

<sup>7</sup> Gradient descend

<sup>8</sup> Chain rule

<sup>9</sup> Local minimum

<sup>10</sup> Learning rate

<sup>11</sup> K-Means clustering algorithm

<sup>12</sup> Epoch

<sup>1</sup> Radial basis

<sup>2</sup> Estimation

<sup>3</sup> Time series prediction

<sup>4</sup> Hidden Layer

<sup>5</sup> Granular activation function having an uncertain mean and fixed standard deviation  $\sigma$

<sup>6</sup> Granular activation function having an uncertain standard deviation  $\sigma$  and fixed mean

انعطاف‌پذیری سیستم را نسبت به آموزش با روش «گرادیان نزولی»، افزایش می‌دهد [۲۳، ۳۷].

از طرفی «الگوریتم خوشه‌بندی K-Means» در شناسایی و پیش‌بینی توابع آشوب با محدودیت‌هایی روبرو است؛ زیرا دو بردار ورودی که فاصله اقلیدسی<sup>۱</sup> کوچکی دارند و ممکن است در لایه پنهان در یک خوشه<sup>۲</sup> قرار گیرند، می‌توانند در این گونه توابع خروجی‌های دور از هم داشته باشند. مگر این که افق پیش‌بینی را در این توابع به قدری کوچک در نظر بگیریم تا دو بردار ورودی که فاصله اقلیدسی کوچکی دارند؛ خروجی‌های نزدیک به هم تولید کنند [۲۷-۲۸]. به همین دلیل؛ بکار بردن «الگوریتم خوشه‌بندی K-Means» سبب می‌گردد که وزن‌های لایه خروجی بار بیشتری را در شناسایی و پیش‌بینی توابع آشوب تحمل نمایند [۳۷]. ولی در مورد شناسایی سیستم‌های غیر خطی پویا که آشوبی نیستند، «الگوریتم خوشه‌بندی K-Means» در تعداد تکرار کمتر عملکرد بهتری دارد.

عملکرد شبکه عصبی RBF گرانولی، بوسیله شناسایی «یک سیستم غیر خطی پویا با پنج ورودی [۶]» و پیش‌بینی «سری زمانی آشوب مکی گلاس [۳۰]» مورد آزمایش قرار می‌گیرد [۲۷]. نویز به کار رفته دارای «نسبت سیگنال به نویز»<sup>۴</sup> برابر صفر، پنج و ده خواهد بود که با شرایط بدون نویز مقایسه می‌گردد [۱۱].

از آنجا که در رابطه با آموزش پارامترهای لایه میانی شبکه عصبی RBF گرانولی، به کمک «الگوریتم خوشه‌بندی K-Means» و «گرادیان نزولی» در شرایط نویزی و بدون نویز و مقایسه آن‌ها با یکدیگر و با دیگر شبکه‌های عصبی و عصبی-فازی؛ پژوهشی انجام نشده است، این مقاله به منظور پاسخ به پرسش‌های زیر شکل گرفته است:

- شبکه عصبی RBF گرانولی، با آموزش پارامترهای لایه میانی بوسیله دو روش «الگوریتم خوشه‌بندی K-Means» یا «گرادیان نزولی»، و پارامترهای لایه خروجی بوسیله روش «گرادیان نزولی» نسبت به شبکه عصبی RBF با آموزش همه پارامترها بوسیله روش «گرادیان نزولی»، در کدام یک از شرایط نویزی و بدون نویز بهتر عمل می‌کند؟

- ایجاد یک ساختار بازه‌ای در دو پارامتر (الف) انحراف استاندارد توابع فعال‌ساز گرانولی که یک پارامتر غیرخطی لایه میانی است و (ب) وزن‌های لایه خروجی که یک پارامتر خطی لایه خروجی است، چه تأثیری در افزایش کارایی سیستم در شرایط نویزی خواهد داشت؟

- در شبکه عصبی RBF گرانولی، آموزش پارامترهای غیرخطی لایه میانی بوسیله روش «الگوریتم خوشه‌بندی K-Means» در مقایسه با «الگوریتم گرادیان نزولی» چه تفاوتی در شناسایی سیستم‌های غیر خطی

پویا و پیش‌بینی توابع آشوب دارد؟ این تفاوت در خطای شناسایی و پیش‌بینی و زمان اجرای برنامه چگونه خواهد بود؟

- کارایی شبکه عصبی RBF گرانولی در مدیریت عدم قطعیت ناشی از شرایط نویزی؛ نسبت به دیگر شبکه‌های عصبی و عصبی-فازی چگونه است؟ آیا به کار بردن شبکه عصبی RBF گرانولی نسبت به دیگر سیستم‌های هوشمند دارای مزیت است؟

## ۲- مرور کارهای انجام شده در زمینه نویز و کاهش اثر آن<sup>۵</sup>

### ۱-۲ تعریف نویز

در یک تعریف کلی، به هر نوسان و تغییر ناخواسته که بر روی سیگنال‌های مورد اندازه‌گیری ظاهر شود، نویز گفته می‌شود. هر کمیتی می‌تواند نویزی گردد [۱۴، ۱۶].

- در مدارهای الکتریکی بیشتر با نویز ولتاژ و جریان سر و کار داریم؛ این نویز ناشی از تغییرات دمایی محیط انتقال انرژی و تأثیر آن بر روی حرکت الکترون‌ها است.

- در حوزه امواج رادیویی و مایکروویو<sup>۶</sup> با نویزهای الکترومغناطیسی و گاهی نیز با نویزی که ناشی از گرما یا تابش<sup>۷</sup> و یون-های کم انرژی باشد روبرو هستیم.

در هر آزمایش دقیق و با کیفیت بالایی که انجام می‌شود؛ باید بتوان نویز محیط را پیش‌بینی و تأثیر آن را کم کرد. اهمیت تحلیل نویز هنگامی کاملاً نمایان می‌شود که متوجه می‌شویم کیفیت سیگنال اندازه‌گیری شده تنها به مقدار انرژی سیگنال بستگی ندارد بلکه به وسیله «نسبت سیگنال به نویز» تعیین می‌شود. نتیجه تحقیقات نشان می‌دهد که بهترین روش برای بهبود «نسبت سیگنال به نویز»، کاهش نویز است نه افزایش قدرت سیگنال [۱۰-۱۱، ۱۶].

### ۲-۲ فرآیند تولید نویز

نویز طبق تعریف، غیر قابل کنترل است و مقدار دقیق آن در آزمایش‌های گوناگون با هم متفاوت است. پس در واقع فرآیند تولید نویز از نوع فرآیندهای تصادفی است و معمولاً تابع توزیع احتمال<sup>۸</sup> [۳۹] متغیر تصادفی نویز را بر اساس قضیه حد مرکزی<sup>۹</sup> [۳۹] به شکل گاوسی یا

<sup>۵</sup> Noise Reduction

<sup>۶</sup> Microwave

<sup>۷</sup> Radiation

<sup>۸</sup> Probability Distribution Function (PDF)

<sup>۹</sup> نظریه حد مرکزی بیان می‌کند که اگر تعداد نمونه‌های یک متغیر تصادفی به

سمت بی‌نهایت میل کند، تابع توزیع احتمال آن متغیر تصادفی به سوی توزیع نرمال میل می‌کند.

<sup>۱</sup> Euclidean distance

<sup>۲</sup> Cluster

<sup>۳</sup> Mackey glass chaotic time series

<sup>۴</sup> Signal to Noise Ratio(SNR)

به ازای هر مقدار  $r$  باید ثابت باشد. برخی از عمومی‌ترین نویزهای موجود، عبارتند از:

• **نویز سفید<sup>۸</sup>**: به یک دنباله نویزی مانند  $(n_r^{(t)})_{r \in N}$  نویز سفید گفته می‌شود اگر متغیر تصادفی  $n_r^{(t)}$  در این دنباله؛ دارای میانگین صفر بوده و داشته باشیم:

$$E((n_r^{(t)})^2) = Var(n_r^{(t)}) = \sigma_n^2 = a \text{ constant value} \quad (4)$$

تابع چگالی توان<sup>۹</sup> و یا طیف توان<sup>۱۰</sup> نویز سفید به فرکانس آن بستگی ندارد و دارای دامنه ثابتی برابر  $\sigma_n^2$  است که به آن توان نویز گفته می‌شود. البته این یک تعریف ایده‌آل است زیرا اگر از یک عدد ثابت نسبت به فرکانس انتگرال بگیریم، واریانس نویز (یا همان توان نویز) بی‌نهایت به دست می‌آید. نویز سفید به دو صورت ظاهر می‌شود؛ نویز دمایی<sup>۱۱</sup> و اثر ساچمه‌ای<sup>۱۲</sup>.

• **آشفته‌گی هارمونیک<sup>۱۳</sup>**: آشفته‌گی‌های هارمونیک در واقع نویزهای تصادفی نیستند بلکه آشفته‌گی‌هایی هستند که از منابع نزدیک بر روی سیستم افتاده است. این نویزها می‌توانند به وسیله طراحی‌های مناسب حذف شوند. روش‌هایی که برای حذف این نویز استفاده می‌شوند عبارتند از؛ پوشش محافظ<sup>۱۴</sup>، زمین کردن<sup>۱۵</sup> مناسب و کاهش حساسیت سیستم<sup>۱۶</sup> به نویز. از آنجا که آشفته‌گی‌های هارمونیک دارای فرکانس-های مشخصی هستند، باعث ایجاد نوسانات نامیرا<sup>۱۷</sup> در سیگنال و ایجاد ضربه<sup>۱۸</sup> در طیف فرکانسی<sup>۱۹</sup> می‌شوند. این رفتار تکین<sup>۲۰</sup> باعث می‌شود که نوع آن‌ها با نویزهای دیگر فرق کند.

• **نویز صورتی<sup>۲۱</sup> یا نویز  $\frac{1}{f}$** : در بررسی سیستم‌ها؛ نویز واقعی سفید نیست بلکه «صورتی» است. به این معنا که دارای فرکانس قطع است. این فرکانس قطع باعث می‌شود که واریانس نویز محدود شود. طیف توان این نویز با آهنگ<sup>۲۲</sup>  $\frac{1}{f}$  کاهش پیدا می‌کند. توان نویز  $\frac{1}{f}$  بستگی به نحوه تولید آن دارد و از وسیله‌ای به وسیله دیگر متفاوت است.

• **نویز آشوبی<sup>۲۳</sup>**: این نویز می‌تواند توسط ماشین‌هایی که دارای قسمت گردنده<sup>۲۴</sup> و بالبه‌های تیز<sup>۲۵</sup> می‌باشند تولید گردد. معمولاً این نویز

نرمال در نظر می‌گیرند. البته در شرایطی که تعداد نمونه‌های «متغیر تصادفی» نویز کم باشد؛ ممکن است توزیع‌های دیگری نیز مد نظر قرار گیرد.

تصادفی بودن نویز سبب می‌شود که تابع توزیع احتمال آن از نوع «نرمال یا گاوسی با میانگین صفر» و به صورت  $N(0, \sigma_n^2)$  در نظر گرفته شود [۹]. بنابراین برای توصیف نویز از مقادیر مربع آن استفاده می‌شود. مقدار مؤثر<sup>۱</sup> نویز از جذر میانگین مربعات<sup>۲</sup> آن به دست می‌آید. البته این پارامتر هیچ اطلاعاتی در مورد چگونگی تغییر مقدار نویز با زمان و یا اجزای فرکانسی آن نمی‌دهد. اگر ویژگی‌های آماری نویز مانند واریانس<sup>۳</sup> یا انحراف استاندارد و یا مقدار مؤثر آن با زمان تغییر نکند به آن نویز ایستا<sup>۴</sup> گفته می‌شود [۹].

در سیستم‌هایی که چند منبع نویز وجود داشته باشد نویز کلی می‌تواند به صورت مجموع نویزهای مختلف در نظر گرفته شود. اگر این نویزها مستقل<sup>۵</sup> از یکدیگر باشند می‌توان مقدار مؤثر را به صورت جمع مقدارهای مؤثر تک تک منابع نویز در نظر گرفت [۳۹-۴۱].

## ۳-۲ انواع نویز

نویزها بیشتر بر اساس تغییرات زمانی و فرکانسی خود از یکدیگر متمایز می‌شوند. در شبکه عصبی بر پایه توابع شعاعی گرانولی، با نویزهایی سروکار داریم که از نوع سیگنال‌های «گسسته در زمان تصادفی» می‌باشند. یک گروه مهم از این سیگنال‌ها، سیگنال‌های WSS<sup>۶</sup> است [۴۳]. یک سیگنال «گسسته در زمان تصادفی» مانند دنباله نویزی  $(n_r^{(t)})_{r \in N}$  را یک سیگنال WSS گویند اگر داشته باشیم:

$$E(n_r^{(t)}) = a \text{ constant value} \quad \forall r \in [-\infty, \infty] \quad (1)$$

$$E(n_p^{(t)} n_q^{(t)}) = E(n_{p+r}^{(t)} n_{q+r}^{(t)}), \quad \forall p, q, r \in N \quad (2)$$

از فرمول‌های (۱) و (۲) نتیجه می‌گیریم که برای WSS بودن یک سیگنال «گسسته در زمان تصادفی» مانند دنباله نویزی  $(n_r^{(t)})_{r \in N}$ ؛ میانگین آن به ازای هر مقدار  $r$  باید ثابت باشد.

$$E((n_r^{(t)})^2) = a \text{ constant value} \Rightarrow Var(n_r^{(t)}) = a \text{ constant value} \quad (3)$$

از فرمول (۳) نتیجه می‌گیریم که برای WSS بودن یک سیگنال «گسسته در زمان تصادفی» مانند دنباله نویزی  $(n_r^{(t)})_{r \in N}$ ؛ واریانس آن

<sup>1</sup> Effective

<sup>2</sup> Root Mean Square (RMS)

<sup>۳</sup> واریانس معرف انرژی نویز است

<sup>4</sup> Static

<sup>۵</sup> نویزهایی مستقل هستند که میانگین حاصل ضرب دو به دوی نویزها صفر شود

<sup>6</sup> Stochastic Discrete Time Signal

<sup>7</sup> Wide-Sense Stationary

<sup>8</sup> White Noise

<sup>9</sup> Power Spectral Density

<sup>10</sup> Power Spectrum

<sup>11</sup> Thermal Noise

<sup>12</sup> Shot Noise

<sup>13</sup> Harmonic Disturbance or Harmonic Oscillation

<sup>14</sup> Shield

<sup>15</sup> Earthing

<sup>16</sup> System Sensitivity Reduction

<sup>17</sup> Undamped Oscillations

<sup>18</sup> Impulse

<sup>19</sup> Frequency Spectrum

<sup>20</sup> Singularity Behavior

<sup>21</sup> Pink Noise or Flicker Noise

<sup>22</sup> Chaotic Noise

<sup>23</sup> Rotating part

<sup>24</sup> Blades

### ۳- معرفی شبکه عصبی RBF

ماتریس  $n$  سطری  $X^{(t)}$  ورودی شبکه عصبی است که به شکل یک دسته داده ورودی به آن وارد می‌شود. و ماتریس  $N^{(t)}$  نویز اضافه شده به این دسته داده ورودی می‌باشد. ماتریس  $w^{j,(t)}$  شامل وزن‌های لایه پنهان شبکه و ماتریس‌های  $c^{j,(t)}$  و  $\sigma^{j,(t)}$  به ترتیب شامل مقادیر مرکز و انحراف استاندارد «توابع فعال‌ساز گاوسی» ترون  $j$  ام لایه پنهان شبکه<sup>۱۲</sup> می‌باشند. اندیس‌های  $r$  و  $t$  به صورت زیر تعریف می‌گردند [۲۳]:

- برای تعداد خوشه‌ها<sup>۱۳</sup> یا تعداد ترون‌های لایه میانی داریم:  
 $j = 1, 2, \dots, m$
- برای ابعاد ورودی‌ها داریم:  $r = 1, 2, \dots, n$
- برای شماره نمونه‌های ورودی داریم:  $t = 1, 2, \dots, T$

همچنین مقادیر اولیه مرکز توابع فعال‌ساز گاوسی ترون‌های لایه پنهان؛ یعنی  $c^{j,(0)}$  معادل  $w^{j,(0)}$  در نظر گرفته می‌شود. تابع  $\psi$  تابع فعال ساز گاوسی ترون‌های لایه میانی است. ترم بکار رفته در این‌جا ترم اقلیدسی می‌باشد. خروجی اولیه ترون  $j$  ام لایه میانی یعنی  $o^{j,(0)}$  را می‌توان به وسیله فرمول زیر به دست آورد [۲۳-۲۶]:

$$o^{j,(0)} = \psi \left[ \sum_{r=1}^n \left( \frac{x_r^{(0)} + n_r^{(0)} - c_r^{j,(0)}}{\sigma_r^{j,(0)}} \right)^2 \right] = \psi \left[ \sum_{r=1}^n \left( u_r^{j,(0)} \right)^2 \right] = \psi \left( d^{j,(0)} \right) = \exp \left( -\frac{1}{2} d^{j,(0)} \right) \quad (5)$$

شکل ماتریسی متغیرهای به کار رفته عبارت است از:

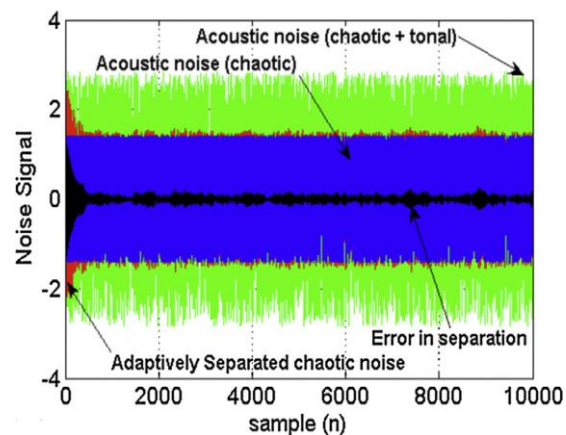
$$X^{(t)} = \begin{bmatrix} x_1^{(t)} \\ x_2^{(t)} \\ \vdots \\ x_n^{(t)} \end{bmatrix}_{n \times 1}, \quad N^{(t)} = \begin{bmatrix} n_1^{(t)} \\ n_2^{(t)} \\ \vdots \\ n_n^{(t)} \end{bmatrix}_{n \times 1}, \quad c^{j,(t)} = \begin{bmatrix} c_1^{j,(t)} \\ c_2^{j,(t)} \\ \vdots \\ c_n^{j,(t)} \end{bmatrix}_{n \times 1} = w^{j,(t)} = \begin{bmatrix} w_1^{j,(t)} \\ w_2^{j,(t)} \\ \vdots \\ w_n^{j,(t)} \end{bmatrix}_{n \times 1}, \quad \sigma^{j,(t)} = \begin{bmatrix} \sigma_1^{j,(t)} \\ \sigma_2^{j,(t)} \\ \vdots \\ \sigma_n^{j,(t)} \end{bmatrix}_{n \times 1} \quad (6)$$

برای مقادیر وزن‌های اتصال بین لایه ورودی و لایه پنهان شبکه؛ ماتریس زیر را خواهیم داشت:

$$W^{(t)} = [w^{1,(t)} \quad w^{2,(t)} \quad \dots \quad w^{m,(t)}]_{1 \times m} = \begin{bmatrix} w_1^{1,(t)} & w_1^{2,(t)} & \dots & w_1^{m,(t)} \\ w_2^{1,(t)} & w_2^{2,(t)} & \dots & w_2^{m,(t)} \\ \vdots & \vdots & \ddots & \vdots \\ w_n^{1,(t)} & w_n^{2,(t)} & \dots & w_n^{m,(t)} \end{bmatrix}_{n \times m} \quad (7)$$

برای مقادیر مرکز توابع فعال‌ساز ترون‌های لایه پنهان؛ ماتریس زیر را خواهیم داشت:

با یک تَن صدا<sup>۱</sup> ترکیب می‌گردد که وابسته به سرعت چرخش قسمت گردان ماشین می‌باشد. بخش آشوبی نویز مربوط به برخورد لبه‌های تیز به هوای اطراف می‌باشد [۹].



شکل ۱. نمودار تغییرات نویز حاصل از ترکیب بخش آشوبی و تَن صدا، پیش و پس از جداسازی این دو بخش، در دامنه زمان [۹]

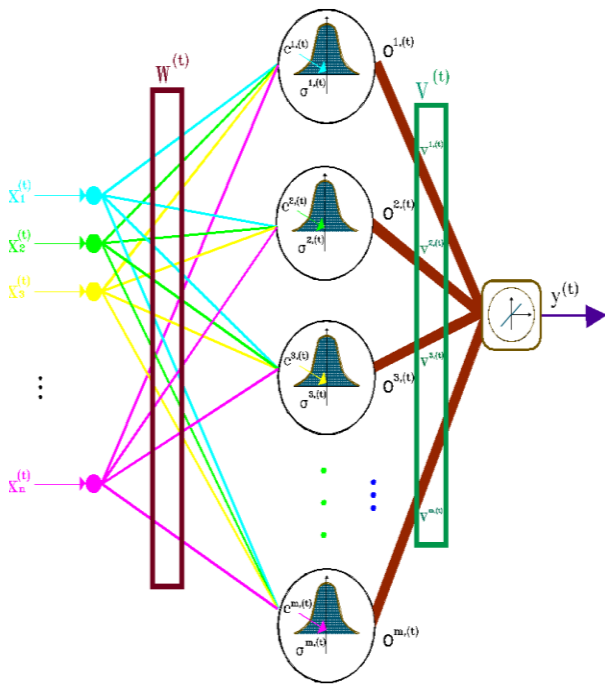
### ۲-۴ کاهش اثر نویز

پدیده نویز در کنترل سیستم‌های مختلف مشکلی عمومی است و سیستم‌های هوشمند بویژه سیستم‌های هوشمند نوع<sup>۲</sup> [۶-۷، ۱۱] که در آن‌ها مدیریت عدم قطعیت به خوبی صورت می‌پذیرد؛ در این حیطه وارد شده‌اند. تشخیص نویز<sup>۳</sup> و حذف نویز<sup>۴</sup> در بسیاری از زمینه‌ها مانند پردازش تصویر در کنار کاهش اثر نویز از اهمیت بالایی برخوردارند. با وجود پیشرفت در طراحی شبکه‌های عصبی، ارائه روش‌هایی که نیاز به دوره آموزش کمتری دارند هنوز مفید است. برخی از این روش‌ها در این‌جا آورده شده است:

- استفاده از فیلتر غیرخطی «کاهش نویز ضربه فازی»<sup>۵</sup> [۱۳-۱۴].
- به کار بردن روش موجک<sup>۶</sup> [۱۲].
- الگوریتم فیلتر خطی FXLMS<sup>۷</sup> [۷].
- الگوریتم کنترل با فیلتر غیرخطی VFSLMS<sup>۸</sup> [۱۵].
- انواع شبکه‌های عصبی مانند شبکه عصبی RBF گرانولی [۱۸، ۲۰] و شبکه عصبی ارتباط تابعی<sup>۹</sup> [۹].
- سیستم‌های منطق فازی نوع<sup>۱۰</sup> [۱۰-۱۱] که به عنوان نمونه می‌توان با شبکه عصبی- فازی نوع<sup>۱۱</sup> [۶-۷] آن را پیاده‌سازی نمود.

<sup>1</sup> Tonal  
<sup>2</sup> Type 2 Intelligent Systems  
<sup>3</sup> Noise Detection  
<sup>4</sup> Noise Canceling  
<sup>5</sup> Fuzzy Impulse Noise Reduction Method (FINRM)  
<sup>6</sup> Wavelet  
<sup>7</sup> Filtered-X Least Mean Square  
<sup>8</sup> Volterra Filtered-X Least Mean Square  
<sup>9</sup> Functiona Link Artificial Neural Network (FLANN)  
<sup>10</sup> Type-2 Fuzzy Logic Systems (T2FLSs)  
<sup>11</sup> Interval Type-2 Fuzzy Neural Network

<sup>12</sup> Gaussian activation functions of the  $j^{\text{th}}$  neuron in the hidden layer of the neural network  
<sup>13</sup> Clusters



شکل ۲. یک شبکه عصبی RBF

### ۱-۳ روش‌های آموزش بکار رفته در شبکه عصبی RBF

برای آموزش پارامترهای غیرخطی لایه میانی که «مرکز» و «انحراف» استاندارد» توابع فعال‌ساز گاوسی می‌باشند می‌توان از دو روش آموزش یکی از نوع آموزش بی‌سرپرست<sup>۱</sup> [۳۳] یعنی «الگوریتم خوشه‌بندی K-Means با m خوشه» [۳۷، ۲۳] و دیگری از نوع آموزش با سرپرست<sup>۲</sup> [۳۳] یعنی «گرادیان نزولی» استفاده نمود [۲۳، ۳۵].

### ۲-۳ الگوریتم خوشه‌بندی K-Means با m خوشه

گام‌های الگوریتم برای آموزش «مرکز توابع فعال‌ساز گاوسی» به شرح زیر است و کاملاً مشابه الگوریتم آموزش «انحراف استاندارد آن‌ها» می‌باشد، جز آنکه در الگوریتم آموزش «انحراف استاندارد» به جای «مرکز توابع فعال‌ساز گاوسی» یعنی  $c_r^{j,(t)}$ ، «انحراف استاندارد» یعنی  $\sigma_r^{j,(t)}$  جایگزین می‌گردد [۳۷، ۲۳]:

۱. به دو پارامتر «مرکز» یا «انحراف استاندارد» توابع فعال‌ساز گاوسی مربوط به نرون‌های لایه میانی، مقادیر تصادفی اولیه به شکل  $c_r^{j,(0)}$  و  $\sigma_r^{j,(0)}$  و در بازه [0 1] نسبت می‌دهیم. برای این دو پارامتر یعنی  $c_r^{j,(t)}$  و  $\sigma_r^{j,(t)}$  اندیس‌های  $j$  و  $r$  همانند گذشته تعریف می‌گردند.

۲. بردار ورودی آموزشی جدید را اعمال می‌کنیم.

$$X^{(t)} + N^{(t)} = [(x_1^{(t)} + n_1^{(t)}) (x_2^{(t)} + n_2^{(t)}) \dots (x_n^{(t)} + n_n^{(t)})] \quad (14)$$

$$C^{(t)} = [c^{1,(t)} \ c^{2,(t)} \ \dots \ c^{m,(t)}]_{1 \times m} = \begin{bmatrix} c_1^{1,(t)} & c_1^{2,(t)} & \dots & c_1^{m,(t)} \\ c_2^{1,(t)} & c_2^{2,(t)} & \dots & c_2^{m,(t)} \\ \vdots & \vdots & \ddots & \vdots \\ c_n^{1,(t)} & c_n^{2,(t)} & \dots & c_n^{m,(t)} \end{bmatrix}_{n \times m} \quad (8)$$

و برای مقادیر انحراف استاندارد این نرون‌ها؛ ماتریس زیر را خواهیم داشت:

$$\Sigma^{(t)} = [\sigma^{1,(t)} \ \sigma^{2,(t)} \ \dots \ \sigma^{m,(t)}]_{1 \times m} = \begin{bmatrix} \sigma_1^{1,(t)} & \sigma_1^{2,(t)} & \dots & \sigma_1^{m,(t)} \\ \sigma_2^{1,(t)} & \sigma_2^{2,(t)} & \dots & \sigma_2^{m,(t)} \\ \vdots & \vdots & \ddots & \vdots \\ \sigma_n^{1,(t)} & \sigma_n^{2,(t)} & \dots & \sigma_n^{m,(t)} \end{bmatrix}_{n \times m} \quad (9)$$

خروجی لایه پنهان شبکه عصبی برای یک دسته داده ورودی برابر خواهد بود با:

$$O^{(t)} = [o^{1,(t)} \ o^{2,(t)} \ \dots \ o^{m,(t)}] \\ = [\psi(X^{(t)}, N^{(t)}, c^{1,(t)}, \sigma^{1,(t)}) \ \psi(X^{(t)}, N^{(t)}, c^{2,(t)}, \sigma^{2,(t)}) \ \dots \ \psi(X^{(t)}, N^{(t)}, c^{m,(t)}, \sigma^{m,(t)})] \quad (10)$$

$$O^{(t)} = \left[ \psi \left[ \sum_{r=1}^n \left( \frac{x_r^{(t)} + n_r^{(t)} - c_r^{1,(t)}}{\sigma_r^{1,(t)}} \right)^2 \right] \ \psi \left[ \sum_{r=1}^n \left( \frac{x_r^{(t)} + n_r^{(t)} - c_r^{2,(t)}}{\sigma_r^{2,(t)}} \right)^2 \right] \ \dots \ \psi \left[ \sum_{r=1}^n \left( \frac{x_r^{(t)} + n_r^{(t)} - c_r^{m,(t)}}{\sigma_r^{m,(t)}} \right)^2 \right] \right] \quad (11)$$

اکنون خروجی نهایی شبکه عصبی برای یک نمونه داده ورودی برابر خواهد بود با [۲۳]:

$$y_N^{(t)} = \sum_{j=1}^m (v^{j,(t)} \times o^{j,(t)}) \quad (12)$$

آنجا که  $v^{j,(t)}$  وزن اتصال بین لایه میانی و لایه خروجی برای  $j$  امین نرون لایه پنهان است و یک پارامتر خطی می‌باشد. در نهایت در یک فرمول کلی برای هر یک از  $t = 1, 2, \dots, T$  نمونه داده ورودی خواهیم داشت [۲۳]:

$$y_N^{(t)} = G(X^{(t)}, N^{(t)}, C^{(t)}, \Sigma^{(t)}, V^{(t)}) = \sum_{j=1}^m [v^{j,(t)} \times \psi(X^{(t)}, N^{(t)}, c^{j,(t)}, \sigma^{j,(t)})] = \sum_{j=1}^m (v^{j,(t)} \times o^{j,(t)}) \quad (13)$$

$$V^{(t)} = \begin{bmatrix} v^{1,(t)} \\ v^{2,(t)} \\ \vdots \\ v^{m,(t)} \end{bmatrix}_{m \times 1} \quad \text{که} \quad \text{و لایه خروجی می‌باشد [۲۳].}$$

<sup>1</sup> Unsupervisory Learning

<sup>2</sup> Supervisory Learning

$$(X^{(t)} + N^{(t)}) \times U^{(t)} = [(x_1^{(t)} + n_1^{(t)}) \cdot u_1^{1(t)} + (x_2^{(t)} + n_2^{(t)}) \cdot u_2^{1(t)} + \dots + (x_n^{(t)} + n_n^{(t)}) \cdot u_n^{1(t)} \\ (x_1^{(t)} + n_1^{(t)}) \cdot u_1^{2(t)} + (x_2^{(t)} + n_2^{(t)}) \cdot u_2^{2(t)} + \dots + (x_n^{(t)} + n_n^{(t)}) \cdot u_n^{2(t)} \dots \\ (x_1^{(t)} + n_1^{(t)}) \cdot u_1^{m(t)} + (x_2^{(t)} + n_2^{(t)}) \cdot u_2^{m(t)} + \dots + (x_n^{(t)} + n_n^{(t)}) \cdot u_n^{m(t)}] \quad (21)$$

هر درایه ماتریس حاصلضرب مربوط به یک خوشه است و هر چه مقدار یک درایه بزرگتر باشد احتمال تعلق بردار ورودی به آن خوشه بیشتر می‌گردد. قطعی‌ترین حالت زمانی است که به جز یک درایه بقیه آن‌ها صفر باشند و بدترین وضعیت زمانی رخ می‌دهد که مقدار دو یا چند درایه مساوی گردد که در این صورت تعلق بردار ورودی به یکی از خوشه‌هایی که مقدار درایه مربوط به آن‌ها مساوی است، به صورت تصادفی تعیین می‌گردد.

۴. تابع هزینه می‌تواند توسط رابطه زیر تعریف گردد:

$$J(t) = \sum_{j=1}^m J_j^{(t)} = \sum_{j=1}^m \|X^{(t)} + N^{(t)} - c^{j,(t)}\|_2 \\ = \sum_{j=1}^m \sqrt{\sum_{r=1}^n (x_r^{(t)} + n_r^{(t)} - c_r^{j,(t)})^2} \quad (22)$$

۵. مرکز خوشه  $c^{j,(t)}$  را به کمک رابطه (۲۳) و انحراف استاندارد خوشه  $\sigma^{j,(t)}$  را به وسیله رابطه (۲۴) تغییر می‌دهیم.

$$c^{j,(t+1)} = c^{j,(t)} + \lambda_c^{(t)}(t)(X^{(t)} + N^{(t)} - c^{j,(t)}) \quad (23)$$

$$\sigma^{j,(t+1)} = \sigma^{j,(t)} + \lambda_\sigma^{(t)}(t)(X^{(t)} + N^{(t)} - \sigma^{j,(t)}) \quad (24)$$

پارامترهای  $\lambda_c^{(t)}$  و  $\lambda_\sigma^{(t)}$  ضرایب آموزش تطبیقی هستند که توسط یک تابع خطی از  $t$  به شکل زیر آموزش می‌بیند.

$$\lambda_c^{(t)} = \lambda_c^{(0)}(1 - \frac{t}{T}) \quad (25)$$

$$\lambda_\sigma^{(t)} = \lambda_\sigma^{(0)}(1 - \frac{t}{T}) \quad (26)$$

پارامترهای  $\lambda_c^{(0)}$  و  $\lambda_\sigma^{(0)}$  ضرایب آموزش تطبیقی اولیه، پارامتر  $t$  دوره آموزشی فعلی و  $T$  تعداد کل تکرار برای آموزش است.

۶. شرط خاتمه حلقه<sup>۲</sup> می‌تواند به یکی از سه شکل زیر در نظر گرفته شود که در صورت برقراری یکی از آن‌ها آموزش پایان می‌یابد و در غیر اینصورت به گام دو برمی‌گردیم.

الف) بعد از طی تعداد تعیین شده‌ای از تکرار.

ب) در صورتی که برای تابع هزینه مقدار معین بدست آمده باشد.

ج) بهبود تابع هزینه نسبت به تکرار قبلی کمتر از یک حد آستانه معین باشد.

۳. به یکی از دو روش زیر، تعلق بردار ورودی آموزشی جدید  $X^{(t)} + N^{(t)}$  را به یک خوشه تعیین می‌کنیم.

۱-۳. نزدیکترین مرکز توابع فعال‌ساز گاوسی؛ نسبت به بردار  $X^{(t)} + N^{(t)}$  را به دست می‌آوریم. در این جا از فاصله اقلیدسی به عنوان معیار نزدیکی بین بردار ورودی  $X^{(t)} + N^{(t)}$  و مرکز خوشه  $c^{j,(t)}$  به صورت زیر استفاده می‌شود.

$$J_j^{(t)} = \sqrt{\|X^{(t)} + N^{(t)} - c^{j,(t)}\|_2} = \sqrt{\sum_{r=1}^n (x_r^{(t)} + n_r^{(t)} - c_r^{j,(t)})^2} \quad (15)$$

۲-۳. تعلق «هر بردار ورودی به یک خوشه» را با استفاده از ماتریس تعلق<sup>۱</sup>  $U$  به ابعاد  $n \times m$  تعیین می‌کنیم. در این ماتریس دودویی، درایه  $u_r^{j,(t)}$  برابر یک است اگر  $r$  آمین داده بردار ورودی یعنی  $x_r^{(t)} + n_r^{(t)}$  به گروه  $j$  تعلق داشته باشد و در غیر این صورت برابر صفر خواهد بود. این ماتریس می‌تواند به صورت زیر توصیف شود:

$$U^{(t)} = \begin{bmatrix} u_1^{1,(t)} & u_1^{2,(t)} & \dots & u_1^{m,(t)} \\ u_2^{1,(t)} & u_2^{2,(t)} & \dots & u_2^{m,(t)} \\ \vdots & \vdots & \ddots & \vdots \\ u_n^{1,(t)} & u_n^{2,(t)} & \dots & u_n^{m,(t)} \end{bmatrix}_{n \times m} \quad (16)$$

که در آن داریم:

$$u_r^{j,(t)} = \begin{cases} 1 & \text{if } \|x_r^{(t)} + n_r^{(t)} - c_r^{j,(t)}\|_2 \leq \|x_r^{(t)} + n_r^{(t)} - c_r^{k,(t)}\|_2 \\ 0 & \text{otherwise} \end{cases}$$

$$\forall k = 1, 2, \dots, m, \forall r = 1, 2, 3, \dots, n \mid k \neq j \quad (17)$$

ماتریس تعلق دودویی  $U$  لازم است که هر دو ویژگی (۱۸) و (۱۹) را داشته باشد.

$$\sum_{j=1}^m u_r^{j,(t)} = 1 \quad \forall r = 1, 2, 3, \dots, n \quad (18)$$

$$\sum_{r=1}^n \sum_{j=1}^m u_r^{j,(t)} = n \quad (19)$$

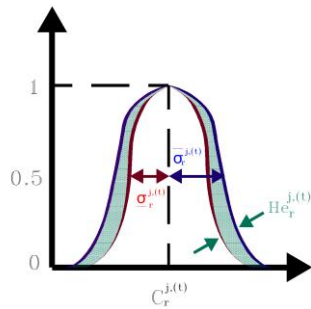
از آنجایی که هر بردار ورودی تنها می‌تواند به یک خوشه تعلق داشته باشد، با ضرب بردار ورودی آموزشی جدید  $X^{(t)} + N^{(t)}$  در ماتریس تعلق  $U$  به ماتریس زیر می‌رسیم:

$$(X^{(t)} + N^{(t)}) \times U^{(t)} = [(x_1^{(t)} + n_1^{(t)}) (x_2^{(t)} + n_2^{(t)}) \dots (x_n^{(t)} + n_n^{(t)})]_{1 \times n} \\ \times \begin{bmatrix} u_1^{1,(t)} & u_1^{2,(t)} & \dots & u_1^{m,(t)} \\ u_2^{1,(t)} & u_2^{2,(t)} & \dots & u_2^{m,(t)} \\ \vdots & \vdots & \ddots & \vdots \\ u_n^{1,(t)} & u_n^{2,(t)} & \dots & u_n^{m,(t)} \end{bmatrix}_{n \times m} \quad (20)$$

<sup>2</sup> End of loop condition

<sup>1</sup> Membership matrix





شکل ۳. تابع فعال‌ساز گاوسی گرانولی در نرون‌های لایه میانی

۱. نمونه‌های دارای بیشترین تکرار را با  $C_r^{j(t)}$  نشان می‌دهند که مرکز یک تابع گاوسی است. در پیاده‌سازی، مراکز توابع گاوسی در یک ماتریس به نام  $(mean^{(t)})^T$  ذخیره می‌گردند. تعداد سطرها این ماتریس همان تعداد ورودی‌های شبکه عصبی و تعداد ستون‌های آن برابر تعداد نرون‌های لایه میانی است. این مقادیر بازه‌ای نیستند [۲۳، ۱۸].

$$(mean^{(t)})^T = C^{(t)} = \begin{bmatrix} C_1^{1(t)} & C_1^{2(t)} & \dots & C_1^{m(t)} \\ C_2^{1(t)} & C_2^{2(t)} & \dots & C_2^{m(t)} \\ \vdots & \vdots & \ddots & \vdots \\ C_n^{1(t)} & C_n^{2(t)} & \dots & C_n^{m(t)} \end{bmatrix}_{1 \times m} = \begin{bmatrix} C_1^{1(t)} & C_1^{2(t)} & \dots & C_1^{m(t)} \\ C_2^{1(t)} & C_2^{2(t)} & \dots & C_2^{m(t)} \\ \vdots & \vdots & \ddots & \vdots \\ C_n^{1(t)} & C_n^{2(t)} & \dots & C_n^{m(t)} \end{bmatrix}_{n \times m} \quad (37)$$

۲. عدم قطعیت اندازه‌گیری شده را با متغیر انتروپی<sup>۳</sup> [۲۳] مدل می‌کنند که در ارتباط با انحراف استاندارد است. انحراف استاندارد دارای مقادیر بازه‌ای است و کران پائین آن در یک ماتریس به نام  $(STDEVleft^{(t)})^T$  و کران بالا در ماتریس دیگری به نام  $(STDEVright^{(t)})^T$  ذخیره می‌گردد. در هر دو ماتریس تعداد سطرها همان تعداد ورودی‌های شبکه عصبی و تعداد ستون‌ها تعداد نرون‌های لایه میانی است [۲۶، ۱۸]. هر چه میزان نویز تزریق شده به داده‌ها بیشتر باشد فاصله کران پائین انحراف استاندارد از کران بالای آن بیشتر می‌گردد و در نتیجه انتروپی بزرگتری خواهیم داشت که با پهن‌تر شدن تابع فعال‌ساز گاوسی گرانولی برای مدیریت عدم قطعیت بزرگتر همراه است [۵، ۴].

$$(STDEVleft^{(t)})^T = \underline{\Sigma}^{(t)} = \begin{bmatrix} \underline{\sigma}_1^{1(t)} & \underline{\sigma}_1^{2(t)} & \dots & \underline{\sigma}_1^{m(t)} \\ \underline{\sigma}_2^{1(t)} & \underline{\sigma}_2^{2(t)} & \dots & \underline{\sigma}_2^{m(t)} \\ \vdots & \vdots & \ddots & \vdots \\ \underline{\sigma}_n^{1(t)} & \underline{\sigma}_n^{2(t)} & \dots & \underline{\sigma}_n^{m(t)} \end{bmatrix}_{1 \times m} = \begin{bmatrix} \underline{\sigma}_1^{1(t)} & \underline{\sigma}_1^{2(t)} & \dots & \underline{\sigma}_1^{m(t)} \\ \underline{\sigma}_2^{1(t)} & \underline{\sigma}_2^{2(t)} & \dots & \underline{\sigma}_2^{m(t)} \\ \vdots & \vdots & \ddots & \vdots \\ \underline{\sigma}_n^{1(t)} & \underline{\sigma}_n^{2(t)} & \dots & \underline{\sigma}_n^{m(t)} \end{bmatrix}_{n \times m} \quad (38)$$

### ۳-۱-۱ الگوریتم پس‌انتشار خطا با استفاده از گرادینان نزولی<sup>۱</sup>

الگوریتم پس انتشار خطا برای  $n$  ورودی و  $m$  خوشه و یک خروجی، بر اساس مجموع مربعات خطا به صورت زیر محاسبه می‌گردد. در این الگوریتم نرخ آموزش می‌تواند برای تمام پارامترها یکسان یا متفاوت در نظر گرفته شود [۳۳-۳۵].

$$e^{(t)} = d^{(t)} - y_N^{(t)}, \quad E^{(t)} = \frac{1}{2} (e^{(t)})^2 \quad (27)$$

برای وزن‌های لایه خروجی که پارامترهای خطی می‌باشند داریم:

$$\Delta v^{j,(t)} = -\eta \times \frac{\partial E^{(t)}}{\partial v^{j,(t)}} = -\eta \times \frac{\partial E^{(t)}}{\partial e^{(t)}} \times \frac{\partial e^{(t)}}{\partial y^{(t)}} \times \frac{\partial y^{(t)}}{\partial v^{j,(t)}} \quad (28)$$

$$\Delta v^{j,(t)} = -\eta \times e^{(t)} \times (-1) \times o^{j,(t)} \quad (29)$$

$$v^{j,(t+1)} = v^{j,(t)} + \Delta v^{j,(t)} = v^{j,(t)} + \eta e^{(t)} o^{j,(t)} \quad (30)$$

برای پارامترهای لایه میانی خواهیم داشت:

$$\Delta c_r^{j,(t)} = -\eta \times \frac{\partial E^{(t)}}{\partial c_r^{j,(t)}} = -\eta \times \frac{\partial E^{(t)}}{\partial e^{(t)}} \times \frac{\partial e^{(t)}}{\partial y^{(t)}} \times \frac{\partial y^{(t)}}{\partial o^{j,(t)}} \times \frac{\partial o^{j,(t)}}{\partial d^{j,(t)}} \times \frac{\partial d^{j,(t)}}{\partial c_r^{j,(t)}} \quad (31)$$

$$\Delta c_r^{j,(t)} = -\eta \times e^{(t)} \times (-1) \times v^{j,(t)} \times \frac{o^{j,(t)} \times d^{j,(t)}}{x_r^{(t)} + n_r^{(t)} - c_r^{j,(t)}} \quad (32)$$

$$c_r^{j,(t+1)} = c_r^{j,(t)} + \Delta c_r^{j,(t)} = c_r^{j,(t)} + \eta e^{(t)} v^{j,(t)} \frac{o^{j,(t)} \times d^{j,(t)}}{x_r^{(t)} + n_r^{(t)} - c_r^{j,(t)}} \quad (33)$$

$$\Delta \sigma_r^{j,(t)} = -\eta \times \frac{\partial E^{(t)}}{\partial \sigma_r^{j,(t)}} = -\eta \times \frac{\partial E^{(t)}}{\partial e^{(t)}} \times \frac{\partial e^{(t)}}{\partial y^{(t)}} \times \frac{\partial y^{(t)}}{\partial o^{j,(t)}} \times \frac{\partial o^{j,(t)}}{\partial d^{j,(t)}} \times \frac{\partial d^{j,(t)}}{\partial \sigma_r^{j,(t)}} \quad (34)$$

$$\Delta \sigma_r^{j,(t)} = -\eta \times e^{(t)} \times (-1) \times v^{j,(t)} \times \frac{o^{j,(t)} \times d^{j,(t)}}{\sigma_r^{j,(t)}} \quad (35)$$

$$\sigma_r^{j,(t+1)} = \sigma_r^{j,(t)} + \Delta \sigma_r^{j,(t)} = \sigma_r^{j,(t)} + \eta e^{(t)} v^j \frac{d^{j,(t)} \times o^{j,(t)}}{\sigma_r^{j,(t)}} \quad (36)$$

### ۴- معرفی شبکه عصبی بر پایه تابع فعال‌ساز

#### گاوسی گرانولی

نمونه‌ای از تابع فعال‌ساز گاوسی گرانولی بکار رفته در نرون‌های لایه میانی، در شکل ۳ نشان داده شده است. به این نوع توابع، تابع گاوسی مدل<sup>۲</sup> ابر<sup>۲</sup> گفته می‌شود و دارای سه مشخصه عددی می‌باشند. این سه مشخصه عددی (۱) مرکز تابع گاوسی (۲) عدم قطعیت اندازه‌گیری شده و (۳) بالاترین مقدار عدم قطعیت هستند که به ترتیب و در ادامه مطلب تعریف شده‌اند [۲۳، ۲۱-۱۷].

<sup>۳</sup> Entropy (En)

<sup>۱</sup> Error Back Propagation with Steepest Gradient Descent

<sup>۲</sup> Cloud model



ماتریس مقادیر ثابت بوده و بازه‌ای نیستند. مقادیر انحراف استاندارد که نامعین می‌باشند و به صورت بازه  $[\underline{\sigma}_r^{j,(t)}, \bar{\sigma}_r^{j,(t)}]$  در نظر گرفته می‌شود نیز در دو ماتریس کران پائین و کران بالا نگهداری می‌گردند.

یک شبکه عصبی RBF گرانولی با استفاده از تابع گاوسی مدل ابر، یک شبکه عصبی گرانولی با مقادیر انحراف استاندارد غیرقطعی<sup>۲</sup> هم نامیده می‌شود [۵]. خروجی هر نرون لایه میانی در این شبکه عصبی RBF گرانولی دارای کران پائین و کران بالا بوده و آن را می‌توان به صورت یک بازه  $[\underline{o}^{j,(t)}, \bar{o}^{j,(t)}]$  نشان داد. ماتریس‌های کران پائین و کران بالای خروجی لایه پنهان شبکه عصبی RBF گرانولی برای یک دسته داده ورودی برابر خواهد بود با [۲۳]:

$$\underline{o}^{(t)} = [\underline{o}^{1,(t)} \quad \underline{o}^{2,(t)} \quad \dots \quad \underline{o}^{m,(t)}] \\ = [\psi(X^{(t)}, \mathcal{N}^{(t)}, c^{1,(t)}, \underline{\sigma}^{1,(t)}) \quad \psi(X^{(t)}, \mathcal{N}^{(t)}, c^{2,(t)}, \underline{\sigma}^{2,(t)}) \quad \dots \quad \psi(X^{(t)}, \mathcal{N}^{(t)}, c^{m,(t)}, \underline{\sigma}^{m,(t)})] \quad (43)$$

$$\underline{o}^{(t)} = \left[ \psi \left[ \sum_{r=1}^n \left( \frac{x_r^{(t)} + n_r^{(t)} - c_r^{1,(t)}}{\underline{\sigma}_r^{1,(t)}} \right)^2 \right] \quad \psi \left[ \sum_{r=1}^n \left( \frac{x_r^{(t)} + n_r^{(t)} - c_r^{2,(t)}}{\underline{\sigma}_r^{2,(t)}} \right)^2 \right] \quad \dots \quad \psi \left[ \sum_{r=1}^n \left( \frac{x_r^{(t)} + n_r^{(t)} - c_r^{m,(t)}}{\underline{\sigma}_r^{m,(t)}} \right)^2 \right] \right] \quad (44)$$

$$\bar{o}^{(t)} = [\bar{o}^{1,(t)} \quad \bar{o}^{2,(t)} \quad \dots \quad \bar{o}^{m,(t)}] \\ = [\psi(X^{(t)}, \mathcal{N}^{(t)}, c^{1,(t)}, \bar{\sigma}^{1,(t)}) \quad \psi(X^{(t)}, \mathcal{N}^{(t)}, c^{2,(t)}, \bar{\sigma}^{2,(t)}) \quad \dots \quad \psi(X^{(t)}, \mathcal{N}^{(t)}, c^{m,(t)}, \bar{\sigma}^{m,(t)})] \quad (45)$$

$$\bar{o}^{(t)} = \left[ \psi \left[ \sum_{r=1}^n \left( \frac{x_r^{(t)} + n_r^{(t)} - c_r^{1,(t)}}{\bar{\sigma}_r^{1,(t)}} \right)^2 \right] \quad \psi \left[ \sum_{r=1}^n \left( \frac{x_r^{(t)} + n_r^{(t)} - c_r^{2,(t)}}{\bar{\sigma}_r^{2,(t)}} \right)^2 \right] \quad \dots \quad \psi \left[ \sum_{r=1}^n \left( \frac{x_r^{(t)} + n_r^{(t)} - c_r^{m,(t)}}{\bar{\sigma}_r^{m,(t)}} \right)^2 \right] \right] \quad (46)$$

در شکل ۴ یک شبکه عصبی RBF گرانولی که بر پایه توابع فعال-ساز گرانولی بنا گردیده، نشان داده شده است. در این شبکه عصبی؛ انحراف استاندارد بازه‌ای در توابع فعال‌ساز گرانولی لایه پنهان و وزن‌های بازه‌ای در لایه خروجی؛ انعطاف‌پذیری و پایداری آن در مقابل نویز را به مقدار زیادی افزایش می‌دهند.

$$(STDEVright^{(t)})^T = \bar{\Sigma}^{(t)} = [\bar{\sigma}^{1,(t)} \quad \bar{\sigma}^{2,(t)} \quad \dots \quad \bar{\sigma}^{m,(t)}]_{1 \times m} = \begin{bmatrix} \bar{\sigma}_1^{1,(t)} & \bar{\sigma}_1^{2,(t)} & \dots & \bar{\sigma}_1^{m,(t)} \\ \bar{\sigma}_2^{1,(t)} & \bar{\sigma}_2^{2,(t)} & \dots & \bar{\sigma}_2^{m,(t)} \\ \vdots & \vdots & \ddots & \vdots \\ \bar{\sigma}_n^{1,(t)} & \bar{\sigma}_n^{2,(t)} & \dots & \bar{\sigma}_n^{m,(t)} \end{bmatrix}_{n \times m} \quad (39)$$

$$Entropy^{(t)} = \bar{\Sigma}^{(t)} - \underline{\Sigma}^{(t)} = [\bar{\sigma}^{1,(t)} \quad \bar{\sigma}^{2,(t)} \quad \dots \quad \bar{\sigma}^{m,(t)}]_{1 \times m} - [\underline{\sigma}^{1,(t)} \quad \underline{\sigma}^{2,(t)} \quad \dots \quad \underline{\sigma}^{m,(t)}]_{1 \times m} \quad (40)$$

$$Entropy^{(t)} = \begin{bmatrix} \bar{\sigma}_1^{1,(t)} - \underline{\sigma}_1^{1,(t)} & \bar{\sigma}_1^{2,(t)} - \underline{\sigma}_1^{2,(t)} & \dots & \bar{\sigma}_1^{m,(t)} - \underline{\sigma}_1^{m,(t)} \\ \bar{\sigma}_2^{1,(t)} - \underline{\sigma}_2^{1,(t)} & \bar{\sigma}_2^{2,(t)} - \underline{\sigma}_2^{2,(t)} & \dots & \bar{\sigma}_2^{m,(t)} - \underline{\sigma}_2^{m,(t)} \\ \vdots & \vdots & \ddots & \vdots \\ \bar{\sigma}_n^{1,(t)} - \underline{\sigma}_n^{1,(t)} & \bar{\sigma}_n^{2,(t)} - \underline{\sigma}_n^{2,(t)} & \dots & \bar{\sigma}_n^{m,(t)} - \underline{\sigma}_n^{m,(t)} \end{bmatrix}_{n \times m} - \begin{bmatrix} \underline{\sigma}_1^{1,(t)} & \underline{\sigma}_1^{2,(t)} & \dots & \underline{\sigma}_1^{m,(t)} \\ \underline{\sigma}_2^{1,(t)} & \underline{\sigma}_2^{2,(t)} & \dots & \underline{\sigma}_2^{m,(t)} \\ \vdots & \vdots & \ddots & \vdots \\ \underline{\sigma}_n^{1,(t)} & \underline{\sigma}_n^{2,(t)} & \dots & \underline{\sigma}_n^{m,(t)} \end{bmatrix}_{n \times m} \quad (41)$$

$$Entropy^{(t)} = \begin{bmatrix} \bar{\sigma}_1^{1,(t)} - \underline{\sigma}_1^{1,(t)} & \bar{\sigma}_1^{2,(t)} - \underline{\sigma}_1^{2,(t)} & \dots & \bar{\sigma}_1^{m,(t)} - \underline{\sigma}_1^{m,(t)} \\ \bar{\sigma}_2^{1,(t)} - \underline{\sigma}_2^{1,(t)} & \bar{\sigma}_2^{2,(t)} - \underline{\sigma}_2^{2,(t)} & \dots & \bar{\sigma}_2^{m,(t)} - \underline{\sigma}_2^{m,(t)} \\ \vdots & \vdots & \ddots & \vdots \\ \bar{\sigma}_n^{1,(t)} - \underline{\sigma}_n^{1,(t)} & \bar{\sigma}_n^{2,(t)} - \underline{\sigma}_n^{2,(t)} & \dots & \bar{\sigma}_n^{m,(t)} - \underline{\sigma}_n^{m,(t)} \end{bmatrix}_{n \times m} \quad (42)$$

۳. بالاترین مقدار عدم قطعیت را با هاپر انتروپی<sup>۱</sup> [۲۳] نشان می‌دهند. هنگامی که کران‌های پائین و بالای انحراف استاندارد یعنی مقادیر  $\underline{\sigma}_r^{j,(t)}$  و  $\bar{\sigma}_r^{j,(t)}$  برای هر تابع فعال‌ساز گرانولی آموزش داده می‌شود، مقدار He تغییر می‌کند. البته مقادیر  $\underline{\sigma}_r^{j,(t)}$  و  $\bar{\sigma}_r^{j,(t)}$  در تابع فعال‌ساز گرانولی، نسبت به خط افقی گذشته از وسط محور عمودی سنجیده می‌شوند. مقدار He طبق شکل ۲ در جایی بوجود می‌آید که تابع فعال‌ساز گرانولی پهن‌ترین موقعیت خود را دارد. منطقی است که با مقدار قویتر نویز در مقادیر ورودی و پس از آموزش کران‌های پائین و بالای انحراف استاندارد، مقدار He بزرگتری بدست آید و برعکس مقدار ضعیف‌تر نویز مقدار He کوچکتری را نتیجه دهد. در واقع بازه‌ای بودن انحراف استاندارد و وزن‌های لایه خروجی، عدم قطعیت حاصل از نویز در داده‌های ورودی را پوشش داده و ابزار غلبه بر آن‌ها را در اختیار شبکه عصبی RBF گرانولی قرار می‌دهد [۶-۱۰].

همانطور که گفته شد، مقادیر مرکز توابع فعال‌ساز گاوسی گرانولی نرون زام لایه میانی، در یک بردار به نام  $c^{j,(t)}$  ذخیره می‌گردد که تعداد درایه‌های آن برابر تعداد ورودی‌های سیستم می‌باشد. درایه‌های این

<sup>2</sup>Uncertain standard deviation

<sup>1</sup>Hyper Entropy (He)

رساندن و جمع کردن مقادیر  $u_r^{j,(t)}$  به ازای تمام ورودی‌های  $x_r^{(t)} + n_r^{(t)}$  با ابعاد  $r = 1, 2, \dots, n$  در ترون  $j$ ام لایه پنهان، مقدار مثبت  $d_r^{j,(t)}$  به دست می‌آید. هنگامی که  $d_r^{j,(t)}$  در ضریب  $-\frac{1}{2}$  ضرب واز حاصل ضرب تابع نمایی به صورت  $\exp\left(-\frac{1}{2}d_r^{j,(t)}\right) = \frac{1}{e^{\frac{1}{2}d_r^{j,(t)}}}$  گرفته می‌شود؛ به دلیل کوچک شدن مخرج کسر، حاصل تقسیم بزرگ شده و کران بالایی خروجی تابع گاوسی گرانولی یعنی  $\bar{o}^{j,(t)}$  که مقداری مثبت است؛ به دست می‌آید.

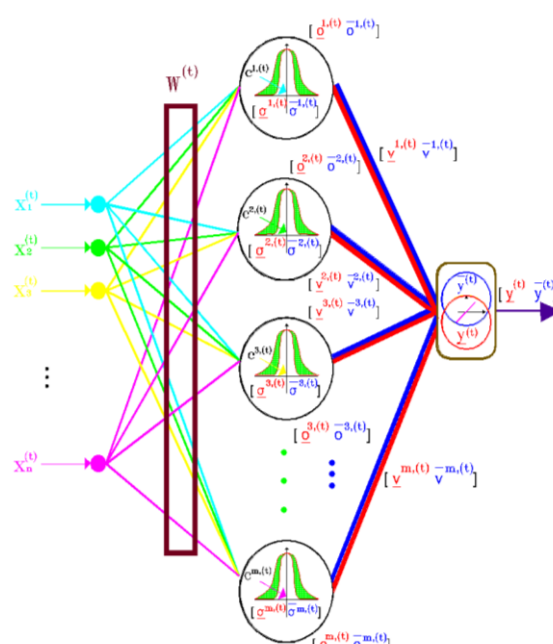
همانطور که مشاهده می‌گردد گاوسی بودن یک تابع فعال‌ساز گرانولی تضمین می‌کند که کران پائین و کران بالایی خروجی آن همیشه مثبت باشد. پس هر مقدار متعلق به بازه  $[\underline{o}^{j,(t)}, \bar{o}^{j,(t)}]$ ، مقداری مثبت خواهد بود.

از طرف دیگر با توجه به این که  $s_r^{j,(t)}$  بازه تغییرات «مرکز وزن» اتصال بین لایه میانی و لایه خروجی برای زائمین ترون لایه پنهان» یعنی  $v_r^{j,(t)}$  می‌باشد. و از آنجا که در محاسبه کران پائین و کران بالایی «وزن» های بازه‌ای در لایه خروجی» یعنی مقادیر  $\bar{v}^{j,(t)}$  و  $\underline{v}^{j,(t)}$  از «قدرمطلق» بازه این وزن‌ها» یعنی  $|s_r^{j,(t)}|$ ، به صورت  $v_r^{j,(t)} = \bar{v}^{j,(t)} - |s_r^{j,(t)}|$  و  $\underline{v}^{j,(t)} = \bar{v}^{j,(t)} + |s_r^{j,(t)}|$  استفاده گردیده است؛ تضمین می‌گردد که به ازای مقادیر منفی و مثبت «مرکز وزن‌های بازه‌ای در لایه خروجی» یعنی  $\bar{v}^{j,(t)}$ ، کران پائین «وزن‌های بازه‌ای در لایه خروجی» یعنی مقدار  $\underline{v}^{j,(t)}$  همیشه کوچکتر از کران بالایی «وزن‌های بازه‌ای در لایه خروجی» یعنی مقدار  $\bar{v}^{j,(t)}$  باشد. یعنی همیشه رابطه نامساوی به صورت  $\underline{v}^{j,(t)} < \bar{v}^{j,(t)}$  که معادل با رابطه نامساوی به شکل  $|s_r^{j,(t)}| < \bar{v}^{j,(t)} - |s_r^{j,(t)}|$  است؛ برقرار می‌باشد.

سرانجام؛ برای محاسبه خروجی نهایی شبکه عصبی RBF گرانولی چهار حالت زیر ممکن است رخ دهد.

- حالت اول: اگر  $v_r^{j,(t)} > 0$  و  $|s_r^{j,(t)}| > \bar{v}^{j,(t)}$  باشد آن‌گاه  $\underline{v}^{j,(t)}, \bar{v}^{j,(t)} > 0$  خواهد بود.
- حالت دوم: اگر  $v_r^{j,(t)} > 0$  و  $|s_r^{j,(t)}| < \bar{v}^{j,(t)}$  باشد آن‌گاه  $\underline{v}^{j,(t)} < 0$  و  $\bar{v}^{j,(t)} > 0$  خواهد بود.
- حالت سوم: اگر  $v_r^{j,(t)} < 0$  و  $|s_r^{j,(t)}| > \bar{v}^{j,(t)}$  باشد آن‌گاه  $\underline{v}^{j,(t)}, \bar{v}^{j,(t)} < 0$  خواهد بود.
- حالت چهارم: اگر  $v_r^{j,(t)} < 0$  و  $|s_r^{j,(t)}| < \bar{v}^{j,(t)}$  باشد آن‌گاه  $\underline{v}^{j,(t)} < 0$  و  $\bar{v}^{j,(t)} > 0$  خواهد بود.

در حالت‌های اول، دوم و چهارم، کران پایین و کران بالایی خروجی نهایی شبکه عصبی RBF گرانولی، به همراه پارامترهای خطی کران پایین و کران بالایی لایه خروجی شبکه به این صورت خواهد بود [۲۳]:



شکل ۴. شبکه عصبی RBF گرانولی بر پایه تابع فعال‌ساز گرانولی

محاسبه کران پائین خروجی یک تابع گاوسی گرانولی در لایه میانی به شکل زیر است [۴، ۵، ۳۷]:

$$\underline{o}^{j,(t)} = \psi \left[ \sum_{r=1}^n \left( \frac{x_r^{(t)} + n_r^{(t)} - c_r^{j,(t)}}{\sigma_r^{j,(t)}} \right)^2 \right] = \psi \left[ \sum_{r=1}^n \left( \frac{\bar{u}_r^{j,(t)}}{\sigma_r^{j,(t)}} \right)^2 \right] = \psi \left( \bar{d}^{j,(t)} \right) = \exp \left( -\frac{1}{2} \bar{d}^{j,(t)} \right) \quad (47)$$

وجود کران پائین انحراف استاندارد؛ به ازای «ورودی  $r$ ام به ترون  $j$ ام لایه پنهان شبکه» یعنی  $\bar{u}_r^{j,(t)}$ ؛ در مخرج کسر  $\frac{x_r^{(t)} + n_r^{(t)} - c_r^{j,(t)}}{\sigma_r^{j,(t)}}$  سبب بزرگ شدن آن و در نتیجه محاسبه  $\bar{u}_r^{j,(t)}$  خواهد شد. با به توان دو رساندن و جمع کردن مقادیر  $\bar{u}_r^{j,(t)}$  به ازای تمام ورودی‌های  $x_r^{(t)} + n_r^{(t)}$  با ابعاد  $r = 1, 2, \dots, n$  در ترون  $j$ ام لایه پنهان، مقدار مثبت  $\bar{d}^{j,(t)}$  به دست می‌آید. هنگامی که  $\bar{d}^{j,(t)}$  در ضریب  $-\frac{1}{2}$  ضرب واز حاصل ضرب تابع نمایی به صورت  $\exp \left( -\frac{1}{2} \bar{d}^{j,(t)} \right) = \frac{1}{e^{\frac{1}{2} \bar{d}^{j,(t)}}}$  گرفته می‌شود؛ به دلیل بزرگ شدن مخرج کسر، حاصل تقسیم کوچک شده و کران پائین خروجی تابع گاوسی گرانولی یعنی  $\underline{o}^{j,(t)}$  که مقداری مثبت است؛ به دست می‌آید.

اکنون برای محاسبه کران بالایی خروجی یک تابع گاوسی گرانولی خواهیم داشت [۴، ۵، ۳۸]:

$$\bar{o}^{j,(t)} = \psi \left[ \sum_{r=1}^n \left( \frac{x_r^{(t)} + n_r^{(t)} - c_r^{j,(t)}}{\sigma_r^{j,(t)}} \right)^2 \right] = \psi \left[ \sum_{r=1}^n \left( \frac{\underline{u}_r^{j,(t)}}{\sigma_r^{j,(t)}} \right)^2 \right] = \psi \left( \underline{d}^{j,(t)} \right) = \exp \left( -\frac{1}{2} \underline{d}^{j,(t)} \right) \quad (48)$$

وجود کران بالایی انحراف استاندارد؛ به ازای «ورودی  $r$ ام به ترون  $j$ ام لایه پنهان شبکه» یعنی  $\underline{u}_r^{j,(t)}$ ؛ در مخرج کسر  $\frac{x_r^{(t)} + n_r^{(t)} - c_r^{j,(t)}}{\sigma_r^{j,(t)}}$  سبب کوچک شدن آن و در نتیجه محاسبه  $\underline{u}_r^{j,(t)}$  خواهد شد. با به توان دو

پائین و کران بالا در «انحراف استاندارد توابع فعال‌ساز گرانولی» به طور جداگانه آموزش می‌بینند. پس به جای دو پارامتر، سه پارامتر برای آموزش وجود دارد [۳۷، ۲۳].

۲-۴ الگوریتم پس‌انتشار خطای گرانولی [۳۳، ۳۵] با استفاده از گرادینان نزولی<sup>۲</sup>

الگوریتم پس‌انتشار خطای گرانولی برای  $n$  ورودی و  $m$  خوشه و یک خروجی، بر اساس روش مجموع مربعات خطا به صورت زیر محاسبه می‌گردد [۴، ۵، ۳۸].

$$e^{(t)} = d^{(t)} - y_N^{(t)}, \quad E^{(t)} = \frac{1}{2} (e^{(t)})^2, \\ e^{(t)} = d^{(t)} - \left( \frac{y_N^{(t)}}{2} + \frac{\bar{y}_N^{(t)}}{2} \right) \quad (55)$$

وزن‌های لایه خروجی را به این صورت آموزش می‌دهیم که برای هر بازه یک «مقدار مرکزی<sup>۳</sup>» و یک «مقدار تغییرات<sup>۴</sup>» که می‌تواند در دو طرف مقدار مرکزی یکسان باشد<sup>۵</sup> و یا یکسان نباشد<sup>۶</sup>، در نظر می‌گیریم. پس دو پارامتر در لایه خروجی آموزش می‌بینند که الف) مرکز وزن‌های بازه‌ای و ب) بازه وزن‌ها می‌باشند. برای مرکز وزن‌های لایه خروجی داریم:

$$\Delta v^{j,(t)} = -\eta \times \frac{\partial E^{(t)}}{\partial v^{j,(t)}} \\ = \left( -\eta \times \frac{\partial E^{(t)}}{\partial e^{(t)}} \times \frac{\partial e^{(t)}}{\partial y^{(t)}} \times \frac{\partial y^{(t)}}{\partial v^{j,(t)}} \right) \\ + \left( -\eta \times \frac{\partial E^{(t)}}{\partial e^{(t)}} \times \frac{\partial e^{(t)}}{\partial \bar{y}^{(t)}} \times \frac{\partial \bar{y}^{(t)}}{\partial v^{j,(t)}} \right) \quad (56)$$

$$\Delta v^{j,(t)} = \left[ -\eta \times e^{(t)} \times \left( -\frac{1}{2} \right) \times \underline{o}^{j,(t)} \right] \\ + \left[ -\eta \times e^{(t)} \times \left( -\frac{1}{2} \right) \times \bar{o}^{j,(t)} \right] \quad (57)$$

$$\Delta v^{j,(t)} = \left( \frac{1}{2} \times \eta \times e^{(t)} \times \underline{o}^{j,(t)} \right) + \left( \frac{1}{2} \times \eta \times e^{(t)} \times \bar{o}^{j,(t)} \right) \quad (58)$$

$$v^{j,(t+1)} = v^{j,(t)} + \Delta v^{j,(t)} = v^{j,(t)} + \eta e^{(t)} \underline{o}^{j,(t)} \quad (59)$$

و برای بازه وزن‌های لایه خروجی داریم:

$$\underline{y}_N^{(t)} = \sum_{j=1}^m (\underline{v}^{j,(t)} \times \underline{o}^{j,(t)}) \\ = \sum_{j=1}^m [(v^{j,(t)} - |s^{j,(t)}|) \times \underline{o}^{j,(t)}] \quad (49)$$

$$\bar{y}_N^{(t)} = \sum_{j=1}^m (\bar{v}^{j,(t)} \times \bar{o}^{j,(t)}) \\ = \sum_{j=1}^m [(v^{j,(t)} + |s^{j,(t)}|) \times \bar{o}^{j,(t)}] \quad (50)$$

در حالت سوم به علت منفی بودن  $\underline{v}^{j,(t)}$  و مثبت بودن  $\bar{v}^{j,(t)}$  و محاسبه کران پایین و کران بالای خروجی نهایی شبکه عصبی RBF گرانولی، به همراه پارامترهای خطی کران پایین و کران بالای لایه خروجی شبکه به این صورت خواهد بود:

$$\underline{y}_N^{(t)} = \sum_{j=1}^m (\underline{v}^{j,(t)} \times \underline{o}^{j,(t)}) \\ = \sum_{j=1}^m [(v^{j,(t)} - |s^{j,(t)}|) \times \underline{o}^{j,(t)}] \quad (51)$$

$$\bar{y}_N^{(t)} = \sum_{j=1}^m (\bar{v}^{j,(t)} \times \bar{o}^{j,(t)}) \\ = \sum_{j=1}^m [(v^{j,(t)} + |s^{j,(t)}|) \times \bar{o}^{j,(t)}] \quad (52)$$

و خروجی نهایی از میانگین کران پائین و کران بالا از رابطه (۵۲) و یا به به فرم تطبیق‌پذیر<sup>۱</sup> از رابطه (۵۳) به دست می‌آید [۶]:

$$y_N^{(t)} = \frac{\underline{y}_N^{(t)} + \bar{y}_N^{(t)}}{2} \quad (53)$$

$$y_N^{(t)} = \alpha_l \times \underline{y}_N^{(t)} + \alpha_u \times \bar{y}_N^{(t)} \quad (54)$$

که در آن  $\alpha_l$  و  $\alpha_u$  به ترتیب پارامترهای تطبیقی کران پائین و کران بالای خروجی هستند.

#### ۱-۴ الگوریتم خوشه‌بندی K-Means با $m$ خوشه

گام‌های الگوریتم برای آموزش پارامترهای لایه میانی شبکه عصبی RBF گرانولی، کاملاً مشابه الگوریتم آموزش برای نوع غیربازه‌ای این شبکه عصبی می‌باشد. تنها تفاوت در این است که در این جا هر دو کران

<sup>2</sup> Granular Error Back Propagation with Steepest Gradient Descent

<sup>3</sup> Midpoint

<sup>4</sup> Interval

<sup>5</sup> Symmetrical

<sup>6</sup> Asymmetrical

<sup>1</sup> Adaptive form

$$\Delta c_r^{j,(t)} = \left[ \frac{1}{2} \times \eta \times e^{(t)} \times \frac{v^{j,(t)} \times \underline{o}^{j,(t)} \times (\bar{u}_r^{j,(t)})^2}{(x_r^{(t)} + n_r^{(t)} - c_r^{j,(t)})} \right] + \left[ \frac{1}{2} \times \eta \times e^{(t)} \times \frac{\bar{v}^{j,(t)} \times \bar{o}^{j,(t)} \times (\underline{u}_r^{j,(t)})^2}{(x_r^{(t)} + n_r^{(t)} - c_r^{j,(t)})} \right] \quad (66)$$

$$\Delta s^{j,(t)} = -\eta \times \frac{\partial E^{(t)}}{\partial s^{j,(t)}} = \left( -\eta \times \frac{\partial E^{(t)}}{\partial e^{(t)}} \times \frac{\partial e^{(t)}}{\partial \underline{y}^{(t)}} \times \frac{\partial \underline{y}^{(t)}}{\partial s^{j,(t)}} \right) + \left( -\eta \times \frac{\partial E^{(t)}}{\partial e^{(t)}} \times \frac{\partial e^{(t)}}{\partial \bar{y}^{(t)}} \times \frac{\partial \bar{y}^{(t)}}{\partial s^{j,(t)}} \right) \quad (60)$$

$$\Delta c_r^{j,(t)} = \left[ \frac{1}{2} \times \eta \times e^{(t)} \times \frac{v^{j,(t)} \times \underline{o}^{j,(t)} \times (\bar{u}_r^{j,(t)})^2}{(x_r^{(t)} + n_r^{(t)} - c_r^{j,(t)})} + \frac{\bar{v}^{j,(t)} \times \bar{o}^{j,(t)} \times (\underline{u}_r^{j,(t)})^2}{(x_r^{(t)} + n_r^{(t)} - c_r^{j,(t)})} \right] \quad (67)$$

$$\Delta s^{j,(t)} = \left[ -\eta \times e^{(t)} \times \left(-\frac{1}{2}\right) \times \underline{o}^{j,(t)} \right] + \left[ -\eta \times e^{(t)} \times \left(-\frac{1}{2}\right) \times \bar{o}^{j,(t)} \right] \quad (61)$$

$$\Delta s^{j,(t)} = \left( \frac{1}{2} \times \eta \times e^{(t)} \times \underline{o}^{j,(t)} \right) + \left( \frac{1}{2} \times \eta \times e^{(t)} \times \bar{o}^{j,(t)} \right) \quad (62)$$

$$s^{j,(t+1)} = s^{j,(t)} + \Delta s^{j,(t)} = s^{j,(t)} + \eta e^{(t)} o^{j,(t)} \quad (63)$$

$$c_r^{j,(t+1)} = c_r^{j,(t)} + \Delta c_r^{j,(t)} = c_r^{j,(t)} + \frac{1}{2} \eta e^{(t)} \left[ \frac{v^{j,(t)} \times \underline{o}^{j,(t)} \times (\bar{u}_r^{j,(t)})^2}{(x_r^{(t)} + n_r^{(t)} - c_r^{j,(t)})} + \frac{\bar{v}^{j,(t)} \times \bar{o}^{j,(t)} \times (\underline{u}_r^{j,(t)})^2}{(x_r^{(t)} + n_r^{(t)} - c_r^{j,(t)})} \right] \quad (68)$$

در لایه میانی سه پارامتر وابسته به تابع فعال‌ساز گرانولی آموزش می‌بینند که الف) مرکز دسته، ب) کران پائین انحراف استاندارد و ج) کران بالای انحراف استاندارد می‌باشند.

برای آموزش کران پائین انحراف استاندارد که از پارامترهای لایه میانی است با فرض  $\bar{o}^{j,(t)} \geq \underline{o}^{j,(t)}$  و  $\bar{d}^{j,(t)} \geq \underline{d}^{j,(t)}$  خواهیم داشت:

برای آموزش مرکز توابع گاوسی که از پارامترهای لایه میانی است و با فرض  $\bar{o}^{j,(t)} \geq \underline{o}^{j,(t)}$  و  $\bar{d}^{j,(t)} \geq \underline{d}^{j,(t)}$  خواهیم داشت [۴، ۵]:

$$\Delta \underline{\sigma}_r^{j,(t)} = -\eta \times \frac{\partial E^{(t)}}{\partial \underline{\sigma}_r^{j,(t)}} = -\eta \times \frac{\partial E^{(t)}}{\partial e^{(t)}} \times \frac{\partial e^{(t)}}{\partial \underline{y}^{(t)}} \times \frac{\partial \underline{y}^{(t)}}{\partial \underline{\sigma}_r^{j,(t)}} \times \frac{\partial \underline{o}^{j,(t)}}{\partial \underline{d}^{j,(t)}} \times \frac{\partial \underline{d}^{j,(t)}}{\partial \underline{\sigma}_r^{j,(t)}} \quad (69)$$

$$\Delta c_r^{j,(t)} = -\eta \times \frac{\partial E^{(t)}}{\partial c_r^{j,(t)}} = \left( -\eta \times \frac{\partial E^{(t)}}{\partial e^{(t)}} \times \frac{\partial e^{(t)}}{\partial \underline{y}^{(t)}} \times \frac{\partial \underline{y}^{(t)}}{\partial \underline{o}^{j,(t)}} \times \frac{\partial \underline{o}^{j,(t)}}{\partial \underline{d}^{j,(t)}} \times \frac{\partial \underline{d}^{j,(t)}}{\partial c_r^{j,(t)}} \right) + \left( -\eta \times \frac{\partial E^{(t)}}{\partial e^{(t)}} \times \frac{\partial e^{(t)}}{\partial \bar{y}^{(t)}} \times \frac{\partial \bar{y}^{(t)}}{\partial \bar{o}^{j,(t)}} \times \frac{\partial \bar{o}^{j,(t)}}{\partial \bar{d}^{j,(t)}} \times \frac{\partial \bar{d}^{j,(t)}}{\partial c_r^{j,(t)}} \right) \quad (64)$$

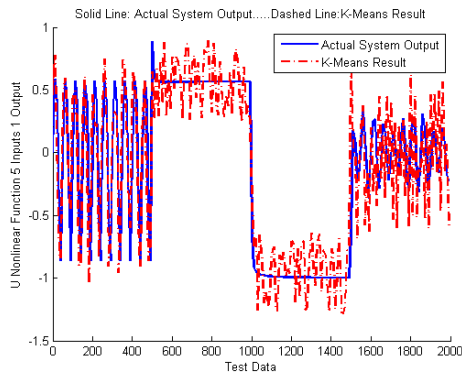
$$\Delta \underline{\sigma}_r^{j,(t)} = -\eta \times e^{(t)} \times \left(-\frac{1}{2}\right) \times (v^{j,(t)} - s^{j,(t)}) \times \left(-\frac{1}{2} \times \underline{o}^{j,(t)}\right) \times \left[-2 \times \frac{(\bar{u}_r^{j,(t)})^2}{\underline{\sigma}_r^{j,(t)}}\right] = \frac{1}{2} \times \eta \times e^{(t)} \times v^{j,(t)} \times \frac{\underline{o}^{j,(t)} \times (\bar{u}_r^{j,(t)})^2}{\underline{\sigma}_r^{j,(t)}} \quad (70)$$

$$\underline{\sigma}_r^{j,(t+1)} = \underline{\sigma}_r^{j,(t)} + \Delta \underline{\sigma}_r^{j,(t)} = \underline{\sigma}_r^{j,(t)} + \frac{1}{2} \eta e^{(t)} v^{j,(t)} \frac{\underline{o}^{j,(t)} \times (\bar{u}_r^{j,(t)})^2}{\underline{\sigma}_r^{j,(t)}} \quad (71)$$

برای آموزش کران بالای انحراف استاندارد که از پارامترهای لایه میانی است با فرض  $\bar{o}^{j,(t)} \geq \underline{o}^{j,(t)}$  و  $\bar{d}^{j,(t)} \geq \underline{d}^{j,(t)}$  خواهیم داشت:

$$\Delta \bar{\sigma}_r^{j,(t)} = -\eta \times \frac{\partial E^{(t)}}{\partial \bar{\sigma}_r^{j,(t)}} = -\eta \times \frac{\partial E^{(t)}}{\partial e^{(t)}} \times \frac{\partial e^{(t)}}{\partial \bar{y}^{(t)}} \times \frac{\partial \bar{y}^{(t)}}{\partial \bar{o}^{j,(t)}} \times \frac{\partial \bar{o}^{j,(t)}}{\partial \bar{d}^{j,(t)}} \times \frac{\partial \bar{d}^{j,(t)}}{\partial \bar{\sigma}_r^{j,(t)}} \quad (72)$$

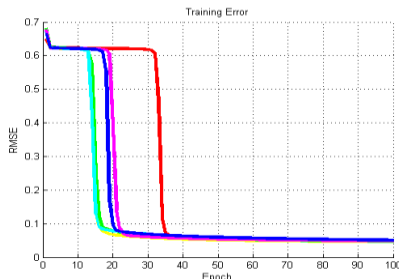
$$\Delta c_r^{j,(t)} = \left[ -\eta \times e^{(t)} \times \left(-\frac{1}{2}\right) \times (v^{j,(t)} - s^{j,(t)}) \times \left(-\frac{1}{2} \times \underline{o}^{j,(t)}\right) \times \left(-2 \times \frac{(\bar{u}_r^{j,(t)})^2}{(x_r^{(t)} + n_r^{(t)} - c_r^{j,(t)})}\right) \right] + \left[ -\eta \times e^{(t)} \times \left(-\frac{1}{2}\right) \times (v^{j,(t)} + s^{j,(t)}) \times \left(-\frac{1}{2} \times \bar{o}^{j,(t)}\right) \times \left(-2 \times \frac{(\underline{u}_r^{j,(t)})^2}{(x_r^{(t)} + n_r^{(t)} - c_r^{j,(t)})}\right) \right] \quad (65)$$



(ب)

شکل ۵. شناسایی «سیستم غیر خطی پویای  $U$  شکل با پنج ورودی» توسط شبکه عصبی RBF گرانولی با ۲۰ نرون لایه میانی و الف) بدون نویز ب)  $SNR=0$

در شکل ۶ خطای آموزش شناسایی «سیستم غیر خطی پویای  $U$  شکل با پنج ورودی» برای ۶ بار اجرای برنامه آورده شده است.



شکل ۶. خطای آموزش شناسایی «سیستم غیر خطی پویای  $U$  شکل با پنج ورودی» توسط شبکه عصبی RBF گرانولی با ۲۰ نرون لایه میانی و  $SNR=0$

تعداد داده‌های به کار رفته برای آموزش در این مثال ۱۵۰۰ داده است. در این مسأله ۴۴۲ داده به عنوان داده آزمایشی به کار رفته‌اند.

در جدول ۱ نتایج شناسایی «سیستم غیر خطی پویای  $U$  شکل با پنج ورودی» توسط شبکه عصبی RBF گرانولی و شبکه عصبی RBF، برای ۱۰۰ بار تکرار الگوریتم آورده شده است.

جدول ۱. نتایج شناسایی «سیستم غیر خطی پویای  $U$  شکل با پنج ورودی» توسط شبکه عصبی RBF گرانولی و شبکه عصبی RBF

| شبکه عصبی RBF       |        | شبکه عصبی RBF گرانولی |        |                           |        | خطا<br>(RMS) | SNR<br>(دسی بل) |
|---------------------|--------|-----------------------|--------|---------------------------|--------|--------------|-----------------|
| گرمادبان نزولی      |        | K-Means               |        | گرمادبان نزولی            |        |              |                 |
| تعداد نرون های لایه |        | تعداد نرون های لایه   |        | تعداد نرون های لایه پنهان |        |              |                 |
| ۲۰                  | ۱۰     | ۲۰                    | ۱۰     | ۲۰                        | ۱۰     | ۱            |                 |
| ۰/۱۷۷۵              | ۰/۱۸۷۶ | ۰/۰۷۰۰                | ۰/۰۷۱۲ | ۰/۰۶۴۹                    | ۰/۰۶۷۲ | ۰/۳۷۸۳       | ۰               |
| ۰/۶۰۴۵              | ۰/۶۳۹۲ | ۰/۴۳۶۵                | ۰/۴۳۷۶ | ۰/۴۳۱۰                    | ۰/۴۳۳۴ | ۰/۵۱۳۴       | آزمون           |
| ۰/۱۵۵۵              | ۰/۱۵۸۴ | ۰/۰۵۵۸                | ۰/۰۶۰۵ | ۰/۰۵۸۰                    | ۰/۰۶۱۹ | ۰/۳۷۸۱       | آموزش           |
| ۰/۴۴۰۰              | ۰/۴۴۴۱ | ۰/۳۵۱۵                | ۰/۳۵۴۲ | ۰/۳۵۲۶                    | ۰/۳۵۵۳ | ۰/۳۸۴۶       | آزمون           |
| ۰/۰۹۷۵              | ۰/۱۰۵۰ | ۰/۰۵۰۰                | ۰/۰۵۵۵ | ۰/۰۵۳۳                    | ۰/۰۵۷۵ | ۰/۲۳۵۰       | آموزش           |
| ۰/۲۸۰۰              | ۰/۳۰۲۸ | ۰/۲۰۳۰                | ۰/۲۰۵۴ | ۰/۲۰۴۱                    | ۰/۲۰۶۹ | ۰/۳۳۹۴       | آزمون           |
| ۰/۰۴۴۵              | ۰/۰۴۷۰ | ۰/۰۴۰۰                | ۰/۰۴۳۲ | ۰/۰۴۶۷                    | ۰/۰۴۹۲ | ۰/۲۲۲۱       | آموزش           |
| ۰/۰۵۲۱              | ۰/۰۵۶۰ | ۰/۰۵۰۵                | ۰/۰۵۵۱ | ۰/۰۵۳۲                    | ۰/۰۵۷۲ | ۰/۳۰۵۶       | آزمون           |
|                     |        |                       |        |                           |        |              | بدون نویز       |

همانطور که از نتایج مندرج در جدول ۱ معلوم می‌گردد؛ در شرایطی که داده‌های ورودی آغشته به نویز نیستند، عملکرد شبکه عصبی RBF

$$\Delta \bar{\sigma}_r^{j,(t)} = -\eta \times e^{(t)} \times \left(-\frac{1}{2}\right) \times (v^{j,(t)} + s^{j,(t)}) \times \left(-\frac{1}{2} \times \bar{\sigma}_r^{j,(t)}\right) \times \left(-2 \times \frac{(y_r^{j,(t)})^2}{\bar{\sigma}_r^{j,(t)}}\right) = \frac{1}{2} \times \eta \times e^{(t)} \times \bar{v}^j \times \frac{\bar{\sigma}_r^{j,(t)} \times \left(\frac{y_r^{j,(t)}}{\bar{\sigma}_r^{j,(t)}}\right)^2}{\bar{\sigma}_r^{j,(t)}} \quad (73)$$

$$\bar{\sigma}_r^{j,(t+1)} = \bar{\sigma}_r^{j,(t)} + \Delta \bar{\sigma}_r^{j,(t)} = \bar{\sigma}_r^{j,(t)} + \frac{1}{2} \eta e^{(t)} \bar{v}^j \times \frac{\bar{\sigma}_r^{j,(t)} \times \left(\frac{y_r^{j,(t)}}{\bar{\sigma}_r^{j,(t)}}\right)^2}{\bar{\sigma}_r^{j,(t)}} \quad (74)$$

## ۵- نتایج شبیه‌سازی

۱-۵ شناسایی سیستم غیر خطی پویای  $U$  شکل با پنج ورودی

هدف شناسایی «سیستم غیر خطی پویای  $U$  شکل با پنج ورودی» است که توسط معادلات زیر تولید می‌گردد [۶].

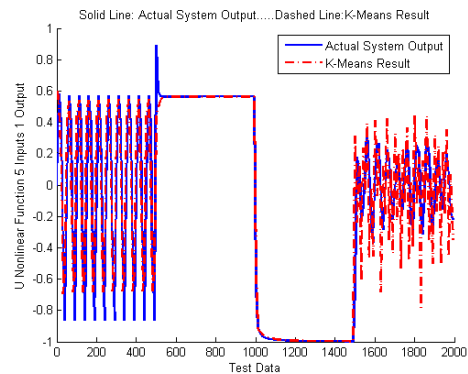
$$y_p(k+1) = f(y_p(k), y_p(k-1), y_p(k-2), u(k), u(k-1)) \quad (75)$$

که در آن:

$$f(x_1, x_2, x_3, x_4, x_5) = \frac{x_1 x_2 x_3 x_5 (x_3 - 1) + x_4}{1 + x_2^2 + x_3^2} \quad (76)$$

$$u(k) = \begin{cases} \sin\left(\frac{\pi k}{25}\right) & 0 < k < 500 \\ +1 & 500 \leq k < 1000 \\ -1 & 1000 \leq k < 1500 \\ 0.3 \sin\left(\frac{\pi k}{25}\right) + 0.1 \sin\left(\frac{\pi k}{32}\right) + 0.6 \sin\left(\frac{\pi k}{10}\right) & 1500 \leq k \leq 2000 \end{cases} \quad (77)$$

در شکل ۵، رنگ «آبی و ممتد» نشان دهنده خروجی واقعی سیستم و رنگ «قرمز و خط چین» نشان دهنده شناسایی سیستم به کمک روش «الگوریتم خوشه‌بندی K-Means» است.



(الف)

گرانولی با آموزش پارامترهای لایه میانی بوسیله روش «الگوریتم خوشه‌بندی K-Means» نسبت به عملکرد همین شبکه با آموزش پارامترهای لایه میانی بوسیله روش «گرادیان نزولی» و شبکه عصبی RBF با آموزش تمام پارامترها بوسیله روش «گرادیان نزولی»، تا حدی بهتر می‌باشد. دلیل این مسأله این است که در شناسایی سیستم‌های غیر خطی پویا و در شرایطی که داده‌های ورودی آغشته به نویز نیستند؛ تعیین تعلق یک دسته داده ورودی به یک خوشه به درستی و به سرعت صورت می‌پذیرد. و این تعیین دقیق خوشه در لایه پنهان؛ مبنای شناسایی درست و سریع در لایه خروجی قرار خواهد گرفت. در حالیکه الگوریتم «گرادیان نزولی» به دلیل احتمال گیر افتادن در مینیمم محلی و دشواری یافتن نرخ آموزش مناسب؛ در تعداد دفعات تکرار مساوی، در شناسایی سیستم‌های غیر خطی پویا ضعیف‌تر عمل می‌کند.

در شرایط بدون نویز، برتری جزئی عملکرد شبکه عصبی RBF نسبت به شبکه عصبی RBF گرانولی با آموزش پارامترهای لایه میانی بوسیله روش «گرادیان نزولی»؛ به دلیل تقریب ناشی از محاسبات بازه‌ای است که روش «گرادیان نزولی» در شبکه عصبی RBF گرانولی را با خطای حاصل از تقریب روبرو می‌سازد.

پس در شرایط بدون نویز، شبکه عصبی RBF گرانولی با آموزش پارامترهای لایه میانی بوسیله روش «الگوریتم خوشه‌بندی K-Means» بهترین عملکرد را دارد و پس از آن شبکه عصبی RBF و در نهایت شبکه عصبی RBF گرانولی با آموزش پارامترهای لایه میانی بوسیله روش «گرادیان نزولی» قرار می‌گیرد.

در شرایطی که داده‌های ورودی آغشته به نویز ضعیف و متوسط هستند نیز برتری روش «الگوریتم خوشه‌بندی K-Means» البته به مقدار کمتر<sup>۱</sup> ادامه می‌یابد. برتری کمتر «الگوریتم خوشه‌بندی K-Means» در حضور نویز به این دلیل است که هر چه دامنه نویز بیشتر می‌گردد؛ تعیین تعلق دسته داده ورودی به یک خوشه مشکل‌تر می‌شود. ولی خطای وجود آمده از این مسأله هنوز آن اندازه بزرگ نشده که بر نقاط ضعف الگوریتم «گرادیان نزولی» غلبه نماید.

هم چنین مشاهده می‌گردد که هر چه میزان نویز تزریق شده به داده‌ها بیشتر باشد؛ عملکرد شبکه عصبی RBF گرانولی با آموزش پارامترهای لایه میانی بوسیله روش «گرادیان نزولی» نسبت به شبکه عصبی RBF بهتر می‌شود. علت این امر این است که خطای حاصل از تقریب در محاسبات بازه‌ای آنقدر یزرگ نیست که سردرگمی شبکه عصبی RBF در مقابل داده‌های ورودی آغشته به نویز را جبران کند. در نتیجه خطای شناسایی شبکه عصبی RBF؛ با افزایش دامنه نویز به سرعت بزرگتر می‌گردد و این شبکه عصبی توانایی شناسایی خود را هر چه بیشتر از دست می‌دهد.

پس در شرایطی که داده‌های ورودی آغشته به نویز ضعیف و متوسط هستند؛ نیز شبکه عصبی RBF گرانولی با آموزش پارامترهای لایه میانی بوسیله روش «الگوریتم خوشه‌بندی K-Means» هم چنان بهترین عملکرد را دارد ولی پس از آن شبکه عصبی RBF گرانولی با آموزش پارامترهای لایه میانی بوسیله روش «گرادیان نزولی» و دست آخر شبکه عصبی RBF قرار می‌گیرد.

نتایج موجود در جدول ۱ همچنین نشان می‌دهد که در نویز با «نسبت سیگنال به نویز» برابر صفر، عملکرد شبکه عصبی RBF گرانولی با آموزش پارامترهای لایه میانی بوسیله روش «گرادیان نزولی» بهتر از حالتی است که پارامترهای لایه میانی بوسیله روش «الگوریتم خوشه‌بندی K-Means» آموزش می‌بیند. چرا که در هنگام وجود نویز قوی در داده ورودی، تعیین تعلق آن به یک خوشه دشوار می‌باشد و خطای وجود آمده در تعیین خوشه بر تقریب ناشی از محاسبات بازه‌ای غلبه می‌یابد.

پس در وضعیت وجود نویز قوی در داده ورودی؛ شبکه عصبی RBF گرانولی با آموزش پارامترهای لایه میانی بوسیله روش «گرادیان نزولی» بهترین عملکرد را دارد و پس از آن شبکه عصبی RBF گرانولی با آموزش پارامترهای لایه میانی بوسیله روش «الگوریتم خوشه‌بندی K-Means» و دست آخر شبکه عصبی RBF قرار می‌گیرد.

جدول ۲ مقایسه نتایج شناسایی «سیستم غیر خطی پویای U شکل با پنج ورودی» بوسیله شبکه عصبی RBF گرانولی، با شبکه عصبی - فازی که آموزش پارامترهای مقدم و تالی در آن به دو روش «بهینه‌سازی گروهی ذرات عمومی»<sup>۲</sup> و «بهینه‌سازی گروهی ذرات بهبود یافته»<sup>۳</sup> صورت پذیرفته را نشان می‌دهد. [۳۶]

جدول ۲. نتایج شناسایی «سیستم غیر خطی پویای U شکل با پنج ورودی» توسط شبکه عصبی RBF گرانولی و شبکه‌های عصبی - فازی دیگر

| روش                                                | شبکه عصبی RBF گرانولی (۴۰ نرون لایه میانی) 100 epochs | شبکه عصبی RBF گرانولی (۳۰ نرون لایه میانی) 100 epochs | خطا (RMS)     |         | SNR (دسی بل)    |
|----------------------------------------------------|-------------------------------------------------------|-------------------------------------------------------|---------------|---------|-----------------|
|                                                    |                                                       |                                                       | گرادیان نزولی | K-Means |                 |
| روش «بهینه‌سازی گروهی ذرات عمومی» 100 epochs       | ۰/۰۳۷۶                                                | ۰/۰۴۱۴                                                | ۰/۰۳۹۱        | ۰/۰۴۴۹  | بدون نویز آموزش |
| روش «بهینه‌سازی گروهی ذرات بهبود یافته» 100 epochs | ۰/۰۴۷۹                                                | ۰/۰۴۸۲                                                | ۰/۰۴۹۷        | ۰/۰۵۰۹  | بدون نویز آزمون |
| روش «بهینه‌سازی گروهی ذرات بهبود یافته» 100 epochs | ؟                                                     | ؟                                                     | ؟             | ؟       | ؟               |

<sup>۲</sup> General Particular Swarm Organization (General PSO)

<sup>۳</sup> Modified Particular Swarm Organization (Modified PSO)

<sup>۱</sup> این مقدار وابسته به دامنه نویز است

دلیل این مسأله این است که در پیش‌بینی توابع آشوب در افق زمانی بزرگ و در شرایطی که داده‌های ورودی آغشته به نویز نیستند؛ تعیین تعلق یک دسته داده ورودی به یک خوشه در لایه پنهان درست و سریع صورت می‌پذیرد. حال از آن جایی که در بحث پیش‌بینی فرض بر آن است که دو بردار ورودی که فاصله اقلیدسی کمی دارند؛ خروجی‌های نزدیک به هم تولید کنند تا سیستم پیش‌بینی‌پذیر گردد. ولی در توابع آشوب در افق زمانی بزرگ، دو بردار ورودی که فاصله اقلیدسی کمی دارند؛ ممکن است که خروجی‌های دور از هم تولید کنند. پس باید پارامترهایی همانند وزن‌های بازه‌ای لایه خروجی؛ درجه آزادی شبکه عصبی RBF گرانولی را به اندازه‌ای بالا ببرند که به کمک آموزش وزن-های بازه‌ای لایه خروجی؛ بتوان این ویژگی توابع آشوب را پوشش داد و تولید خروجی‌های دور از هم را پیش‌بینی نمود. در نهایت در نبود نقاط ضعف روش «گرادیان نزولی» یعنی احتمال گیر افتادن در مینیمم محلی و دشواری یافتن نرخ آموزش مناسب؛ برتری در شرایط فاقد نویز با روش «الگوریتم خوشه‌بندی K-Means» خواهد بود.

سرانجام در شبکه عصبی RBF، به علت بازه‌ای نبودن وزن‌های لایه خروجی؛ انعطاف‌پذیری شبکه عصبی برای پیش‌بینی تولید خروجی‌های دور از هم در توابع آشوب؛ کمتر شده و همین مسأله وجود خطای ناشی از تقریب در محاسبات بازه‌ای در شبکه عصبی RBF گرانولی را پوشش داده و دو شبکه عصبی RBF گرانولی با آموزش پارامترهای لایه میانی بوسیله روش «گرادیان نزولی» و شبکه عصبی RBF با آموزش تمام پارامترها بوسیله روش «گرادیان نزولی» را به نتایج مشابهی می‌رساند. در حالی که شبکه عصبی RBF گرانولی با آموزش پارامترهای لایه میانی بوسیله روش «الگوریتم خوشه‌بندی K-Means» به علت رهایی از روش «گرادیان نزولی» در آموزش پارامترهای لایه پنهان؛ به نتایج بهتری دست پیدا می‌کند.

پس در شرایط بدون نویز، شبکه عصبی RBF گرانولی با آموزش پارامترهای لایه میانی بوسیله روش «الگوریتم خوشه‌بندی K-Means» بهترین عملکرد را دارد و پس از آن شبکه عصبی RBF و شبکه عصبی RBF گرانولی با آموزش پارامترهای لایه میانی بوسیله روش «گرادیان نزولی» در یک رده قرار می‌گیرند.

علت برتری روش «گرادیان نزولی» در شبکه عصبی RBF گرانولی؛ و در هنگام شناسایی و پیش‌بینی توابع آشوب در افق زمانی بزرگ؛ به ویژه در حضور نویز با «نسبت سیگنال به نویز» برابر صفر و پنج و تا حد کمتری ده، این است که چون در آموزش پارامترهای لایه میانی بوسیله روش «الگوریتم خوشه‌بندی K-Means»، و با حضور نویز متوسط و قوی و تا حدی نویز ضعیف، تعیین تعلق دسته داده ورودی به یک خوشه به درستی و به سرعت صورت نگرفته است؛ پس آموزش وزن‌های بازه‌ای لایه خروجی در شناسایی و پیش‌بینی درست توابع آشوب کارساز نیست

خطای آزمون برای این شبکه عصبی- فازی در دسترس نبود. ولی از مقایسه خطای آموزش می‌توان پی برد که روند شناسایی «سیستم غیر خطی پویای U شکل با پنج ورودی» بوسیله شبکه عصبی RBF گرانولی فارغ از نوع آموزش پارامترهای آن؛ نسبت به شبکه عصبی- فازی با آموزش پارامترهای مقدم و تالی به روش «بهینه‌سازی گروهی ذرات عمومی» بهتر صورت می‌پذیرد.

هر چند که در ۱۰۰ بار تکرار الگوریتم، شبکه عصبی- فازی با آموزش پارامترهای مقدم و تالی به روش «بهینه‌سازی گروهی ذرات بهبود یافته» توانایی بیشتری را در شناسایی «سیستم غیر خطی پویای U شکل با پنج ورودی» نشان می‌دهد. ولی به علت زیاد بودن تعداد پارامترها و کندی الگوریتم‌های تکاملی از جمله الگوریتم تکاملی «بهینه‌سازی گروهی ذرات بهبود یافته»، به نظر می‌رسد که این الگوریتم‌ها باید زمان اجرای بزرگتری داشته باشند.

## ۲-۵ پیش‌بینی سری زمانی مکی گلاس با توجه به وجود نویز ورودی

این سری زمانی توسط معادله زیر تولید می‌گردد [۳۰].

$$\dot{x} = \frac{0.2x(t-\tau)}{1+x^{10}(t-\tau)} - 0.1x(t) \quad (78)$$

از ۱۲۰۱ داده تولید شده به وسیله این سری زمانی؛ تعداد ۸۴۰ داده به عنوان داده آموزشی<sup>۱</sup> و ۳۵۲ تا به عنوان داده آزمایشی<sup>۲</sup> به کار رفته‌اند.

جدول ۳. نتایج پیش‌بینی سری زمانی مکی گلاس توسط شبکه عصبی RBF

گرانولی و شبکه عصبی RBF

| خطا<br>(RMS) | SNR<br>(دسی بل) | شبکه عصبی RBF گرانولی     |                           |                           |                           |                           |                           |        |    |
|--------------|-----------------|---------------------------|---------------------------|---------------------------|---------------------------|---------------------------|---------------------------|--------|----|
|              |                 | گرادیان نزولی             |                           |                           | K-Means                   |                           |                           |        |    |
|              |                 | تعداد نرون‌های لایه پنهان | تعداد نرون‌های لایه میانی | تعداد نرون‌های لایه خروجی | تعداد نرون‌های لایه پنهان | تعداد نرون‌های لایه میانی | تعداد نرون‌های لایه خروجی |        |    |
| ۰            | ۵               | ۲۰                        | ۱۰                        | ۲۰                        | ۱۰                        | ۲۰                        | ۱۰                        | ۲۰     | ۱۰ |
| آموزش        |                 | ۰/۰۲۲۸                    | ۰/۰۲۰۷                    | ۰/۰۴۵۲                    | ۰/۰۴۰۰                    | ۰/۰۱۳۲                    | ۰/۰۱۰۸                    |        |    |
| آزمون        |                 | ۰/۲۱۱۱                    | ۰/۲۱۰۹                    | ۰/۳۲۰۸                    | ۰/۳۲۳۶                    | ۰/۳۲۲۵                    | ۰/۳۲۳۶                    |        |    |
| ۵            | ۱۰              | ۲۰                        | ۱۰                        | ۲۰                        | ۱۰                        | ۲۰                        | ۱۰                        | ۲۰     | ۱۰ |
| آموزش        |                 | ۰/۰۳۸۲                    | ۰/۰۲۰۶                    | ۰/۰۱۵۷                    | ۰/۰۳۴۵                    | ۰/۰۳۱۵                    | ۰/۰۵۶۶                    | ۰/۰۵۴۵ |    |
| آزمون        |                 | ۰/۱۱۸۳                    | ۰/۱۱۹۰                    | ۰/۱۱۷۰                    | ۰/۲۴۰۰                    | ۰/۲۲۷۵                    | ۰/۴۴۷۴                    | ۰/۳۰۵۷ |    |
| ۱۰           | بدون<br>هزنه    | ۲۰                        | ۱۰                        | ۲۰                        | ۱۰                        | ۲۰                        | ۱۰                        | ۲۰     | ۱۰ |
| آموزش        |                 | ۰/۰۳۴۵                    | ۰/۰۱۸۳                    | ۰/۰۱۳۸                    | ۰/۰۲۹۳                    | ۰/۰۲۷۵                    | ۰/۰۳۱۷                    | ۰/۰۳۰۴ |    |
| آزمون        |                 | ۰/۰۶۹۱                    | ۰/۰۶۶۸                    | ۰/۰۶۶۲                    | ۰/۰۹۹۵                    | ۰/۰۷۸۰                    | ۰/۱۱۲۸                    | ۰/۱۰۰۲ |    |
| بدون<br>هزنه | بدون<br>هزنه    | ۲۰                        | ۱۰                        | ۲۰                        | ۱۰                        | ۲۰                        | ۱۰                        | ۲۰     | ۱۰ |
| آموزش        |                 | ۰/۰۳۱۰                    | ۰/۰۱۱۳                    | ۰/۰۱۱۱                    | ۰/۰۱۱۱                    | ۰/۰۱۱۰                    | ۰/۰۱۱۳                    | ۰/۰۱۱۰ |    |
| آزمون        |                 | ۰/۰۳۲۴                    | ۰/۰۱۳۷                    | ۰/۰۱۲۸                    | ۰/۰۱۳۱                    | ۰/۰۱۲۷                    | ۰/۰۱۳۷                    | ۰/۰۱۲۷ |    |

همانطور که از نتایج مندرج در جدول ۳ معلوم می‌گردد؛ در شرایطی که داده‌های ورودی آغشته به نویز نیستند، عملکرد شبکه عصبی RBF گرانولی با آموزش پارامترهای لایه میانی بوسیله روش «الگوریتم خوشه‌بندی K-Means» نسبت به عملکرد همین شبکه با آموزش پارامترهای لایه میانی بوسیله روش «گرادیان نزولی» و شبکه عصبی RBF با آموزش تمام پارامترها بوسیله روش «گرادیان نزولی»، اندکی بهتر می‌باشد.

<sup>1</sup> Training data set

<sup>2</sup> Test data set



در جدول ۴ مقایسه خطای آزمون پیش‌بینی «سری زمانی مکی-گلاس»؛ توسط «شبکه عصبی RBF گرانولی» با «سیستم‌های منطق فازی نوع ۱ و نوع ۲» [۱۱] آورده شده است.

جدول ۴. نتایج پیش‌بینی سری زمانی مکی گلاس توسط شبکه عصبی RBF گرانولی و سیستم‌های منطق فازی نوع ۱ و نوع ۲

|        |        |                                      |
|--------|--------|--------------------------------------|
| ۰      | ۱۰     | SNR (دسی بل)                         |
| آزمون  | آزمون  | خطا (RMS)                            |
| ۰/۳۱۰۷ | ۰/۰۶۶۲ | شبکه عصبی RBF<br>(تست لایه میانی)    |
| ۰/۳۳۳۶ | ۰/۰۷۸۰ |                                      |
| ۰/۱۱۰۵ | ۰/۰۱۲۱ | شبکه عصبی RBF<br>(دوینست لایه میانی) |
| ۰/۳۰۱۵ | ۰/۰۱۱۸ |                                      |
| ۰/۱۴۲۹ | ۰/۰۸۳۸ | سیستم های منطقی فازی                 |
| ۰/۱۴۲۲ | ۰/۰۹۳۵ |                                      |
| ۰/۱۵۱۷ | ۰/۰۹۵۳ |                                      |

همان‌طور که از مقایسه نتایج مندرج در جدول ۴ مشخص است؛ در حضور نویز ضعیف عملکرد «شبکه عصبی RBF گرانولی» با هر دو روش آموزش؛ بهتر از «سیستم‌های منطق فازی نوع ۱ و نوع ۲» می‌باشد. که با توجه به بیشتر بودن تعداد پارامترهای آموزش دیده در «سیستم‌های منطق فازی نوع ۱ و بویژه نوع ۲» نسبت به «شبکه عصبی RBF گرانولی»، و پیچیدگی بیشتر الگوریتم آموزشی آن‌ها، زمان اجرای برنامه در «سیستم‌های منطق فازی نوع ۱ و نوع ۲» به شدت افزایش می‌یابد. و این موضوع برتری «شبکه عصبی RBF گرانولی» را برجسته‌تر می‌سازد.

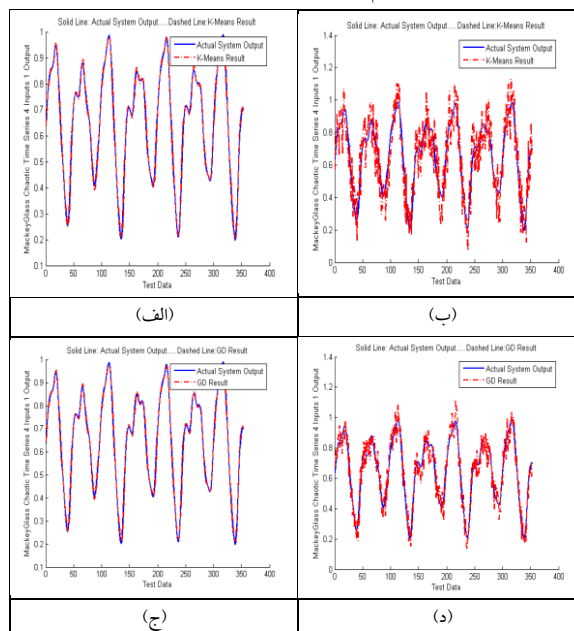
در حضور نویز قوی، عملکرد «شبکه عصبی RBF گرانولی» با افزایش تعداد ترون‌های لایه میانی؛ می‌تواند نسبت به عملکرد «سیستم‌های منطق فازی نوع ۱ و نوع ۲» برتری یابد. زیرا با افزایش تعداد ترون‌های لایه میانی؛ درجه آزادی سیستم با افزایش تعداد پارامترهای آموزش دیده؛ زیاد می‌گردد. به عنوان مثال با داشتن ۳۰۰ ترون در لایه پنهان، خطای آزمون برای روش «الگوریتم خوشه‌بندی K-Means» به ۰/۱۵ کاهش می‌یابد که از «سیستم‌های منطق فازی نوع ۱» بهتر است. البته زمان اجرای برنامه به مقدار زیادی افزایش می‌یابد ولی در مقایسه با زمان اجرای چند دقیقه‌ای «سیستم‌های منطق فازی نوع ۲»، زمان اجرای کمتر از یک دقیقه (جدول ۷) در شبکه عصبی RBF گرانولی؛ مناسب به نظر می‌رسد.

و زمان کافی برای آموزش این وزن‌ها در تعداد تکرار محدود الگوریتم وجود ندارد. از اینرو می‌توان گفت که خطای بوجود آمده در تعیین خوشه در «الگوریتم خوشه‌بندی K-Means»، آنقدر بزرگ است که بر خطای تقریب ناشی از محاسبات بازه‌ای، در «الگوریتم گرادینان نزولی» برتری می‌یابد و روش «گرادینان نزولی» بهتر عمل می‌نماید.

سرانجام در مورد شبکه عصبی RBF، می‌توان گفت که در مقابل داده‌های ورودی آغشته به نویز؛ دچار سردرگمی در شناسایی و پیش‌بینی توابع آشوب می‌گردد. زیرا به علت گرانولی نبودن توابع فعال‌ساز در لایه پنهان و بازه‌ای نبودن وزن‌های لایه خروجی، انعطاف‌پذیری کافی در مواجهه با نویز را ندارد.

پس در وضعیت وجود نویز در داده ورودی؛ شبکه عصبی RBF گرانولی با آموزش پارامترهای لایه میانی بوسیله روش «گرادینان نزولی» بهترین عملکرد را دارد و پس از آن شبکه عصبی RBF گرانولی با آموزش پارامترهای لایه میانی بوسیله روش «الگوریتم خوشه‌بندی K-Means» و دست آخر شبکه عصبی RBF قرار می‌گیرد.

در شکل ۷؛ رنگ «آبی و ممتد» نشان دهنده خروجی واقعی سیستم و رنگ «قرمز و خط چین» نشان دهنده پیش‌بینی سری زمانی مکی گلاس توسط شبکه عصبی RBF گرانولی به کمک «الگوریتم خوشه‌بندی K-Means» و «الگوریتم گرادینان نزولی» است. همان‌طور که انتظار می‌رود در وضعیت عدم وجود نویز در داده ورودی؛ «الگوریتم خوشه‌بندی K-Means» عملکرد بهتری دارد در حالی که در حضور نویز قوی در داده ورودی؛ عملکرد «الگوریتم گرادینان نزولی» بهتر می‌گردد.

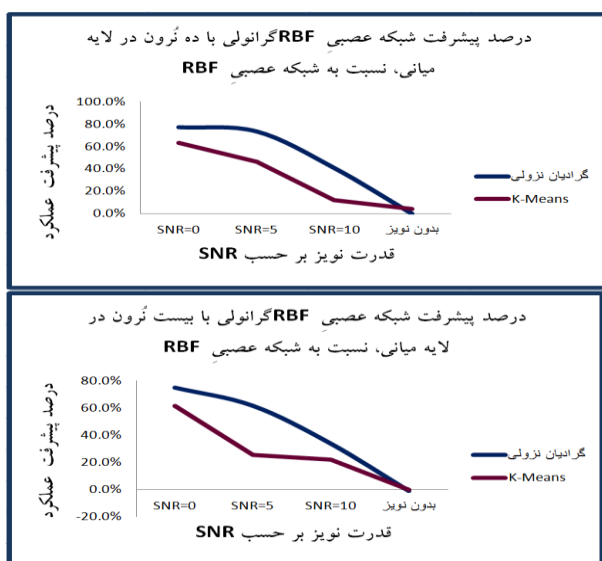


شکل ۷. پیش‌بینی سری زمانی مکی گلاس توسط شبکه عصبی RBF گرانولی با ۲۰ ترون لایه میانی و به روش الف) «الگوریتم خوشه‌بندی K-Means» و بدون نویز ب) «الگوریتم خوشه‌بندی K-Means» و SNR=0 ج) «الگوریتم گرادینان نزولی» و بدون نویز د) «الگوریتم گرادینان نزولی» و SNR=-10

جدول ۶. درصد بهبود عملکرد شبکه عصبی RBF گرانولی نسبت به شبکه عصبی RBF در پیش‌بینی سری زمانی آشوب مکی گلاس

| خطای آزمون (بیست ترون لایه میانی) | خطای آزمون (ده ترون لایه میانی) |               | خطای آزمون (ده ترون لایه میانی) | خطای آزمون (بیست ترون لایه میانی) |               |
|-----------------------------------|---------------------------------|---------------|---------------------------------|-----------------------------------|---------------|
|                                   | شبکه عصبی گرانولی               | شبکه عصبی RBF | شبکه عصبی گرانولی               | شبکه عصبی RBF                     | شبکه عصبی RBF |
| SNR (دسی بل)                      | گرایان نزولی                    | K-Means       | گرایان نزولی                    | K-Means                           | گرایان نزولی  |
| ۰                                 | ۰/۲۱۰۷                          | ۰/۳۲۳۶        | ۰/۱۴۳۶                          | ۰/۳۴۰۸                            | ۰/۹۲۲۵        |
| ۵                                 | ۰/۱۱۷۰                          | ۰/۲۲۷۵        | ۰/۳۰۵۷                          | ۰/۲۴۰۰                            | ۰/۴۴۷۴        |
| ۱۰                                | ۰/۰۶۶۲                          | ۰/۰۷۸۰        | ۰/۱۰۰۲                          | ۰/۰۹۹۵                            | ۰/۱۱۲۸        |
| بدون نویز                         | ۰/۰۱۲۸                          | ۰/۰۱۲۷        | ۰/۰۱۲۷                          | ۰/۰۱۳۱                            | ۰/۰۱۳۷        |
| درصد بهبود عملکرد                 | درصد بهبود                      |               | درصد بهبود                      |                                   | درصد بهبود    |
|                                   | گرایان نزولی                    | K-Means       | گرایان نزولی                    | K-Means                           |               |
| ۰                                 | ۰/۷۵/۰                          | ۰/۶۱/۶        | ۰/۷۷/۱                          | ۰/۶۳/۱                            | ۰/۷۷/۱        |
| ۵                                 | ۰/۶۱/۷                          | ۰/۲۵/۶        | ۰/۷۳/۴                          | ۰/۴۶/۴                            | ۰/۷۳/۴        |
| ۱۰                                | ۰/۳۳/۹                          | ۰/۲۲/۲        | ۰/۴۰/۸                          | ۰/۱۱/۸                            | ۰/۴۰/۸        |
| بدون نویز                         | ۰/۰/۸                           | ۰/۰/۰         | ۰/۰/۰                           | ۰/۰/۰                             | ۰/۰/۰         |

در شکل ۸ که مربوط به پیش‌بینی «سری زمانی آشوب مکی گلاس» می‌باشد، نمودار کاهش درصد پیشرفت عملکرد شبکه عصبی RBF گرانولی نسبت به شبکه عصبی RBF در هنگام کاهش میزان نویز تزریق شده به داده‌های ورودی و برای هر دو روش آموزشی؛ نشان داده شده است.



شکل ۸. عملکرد شبکه عصبی RBF گرانولی نسبت به شبکه عصبی RBF

در جدول ۵ مقایسه خطای آموزش پیش‌بینی «سری زمانی مکی-گلاس»؛ توسط «شبکه عصبی RBF گرانولی» با «شبکه عصبی RBF قوی با هرس و رشد ترون‌های لایه میانی [۱۶]» آورده شده است.

جدول ۵. نتایج پیش‌بینی سری زمانی مکی گلاس توسط شبکه عصبی RBF گرانولی و شبکه عصبی RBF قوی با هرس و رشد ترون‌های لایه میانی

| خطا (RMS) | شبکه عصبی RBF (ده ترون لایه میانی) 100 epochs |         | شبکه عصبی RBF (بیست ترون لایه میانی) 100 epochs |         | شبکه عصبی RBF 100 epochs |
|-----------|-----------------------------------------------|---------|-------------------------------------------------|---------|--------------------------|
|           | گرایان نزولی                                  | K-Means | گرایان نزولی                                    | K-Means |                          |
| آموزش     | ۰/۰۲۰۶                                        | ۰/۰۳۴۵  | ۰/۰۱۵۷                                          | ۰/۰۳۱۵  | ۰/۰۵۴۵                   |
| آزمون     | ۰/۰۱۸۳                                        | ۰/۰۲۹۳  | ۰/۰۱۳۸                                          | ۰/۰۲۷۵  | ۰/۰۳۰۴                   |

خطای آزمون برای این شبکه عصبی موجود نبود ولی از مقایسه خطای آموزش می‌توان پی برد که روند پیش‌بینی «سری زمانی مکی-گلاس» بوسیله شبکه عصبی RBF گرانولی فارغ از نوع آموزش پارامترهای آن؛ نسبت به «شبکه عصبی RBF قوی با هرس و رشد ترون-های لایه میانی» بهتر صورت می‌پذیرد. و این برتری با قویتر شدن نویز، نمود بیشتری می‌یابد. علت این امر وجود تابع فعال‌ساز گرانولی در ساختار شبکه عصبی RBF گرانولی است که با تنظیم «پارامترهای مربوط به پهنای عدم قطعیت» یعنی کران پائین و کران بالای انحراف استاندارد؛ مدیریت عدم قطعیت در شبکه عصبی RBF گرانولی را انجام می‌دهد.

در جدول ۶ درصد پیشرفت عملکرد شبکه عصبی RBF گرانولی نسبت به شبکه عصبی RBF در هنگام کاهش میزان نویز تزریق شده به داده‌ها آورده شده است. در شبکه عصبی RBF گرانولی، هر دو روش «گرایان نزولی» و «الگوریتم خوشه‌بندی K-Means» برای آموزش پارامترهای لایه میانی مد نظر قرار گرفته است. خطای آزمون مربوط به پیش‌بینی «سری زمانی آشوب مکی گلاس» می‌باشد و در توابع آشوب دیگر مانند «شناسایی جاذب لورنز»<sup>۱</sup> [۳۲] و «شناسایی نگاشت آشوبگونه ایکدا»<sup>۲</sup> و «شناسایی نگاشت هنون»<sup>۳</sup> [۳۱] نیز انتظار این است که همین رفتار مشاهده گردد [۲۷-۳۲].

<sup>۱</sup> Lorenz attractor

<sup>۲</sup> Daisaku Ikeda

<sup>۳</sup> Henon map

## ۳-۵ عوامل مؤثر بر زمان اجرای برنامه

همان‌طور که از نتایج به دست آمده از شناسایی سیستم غیر خطی پویای  $U$  شکل با پنج ورودی و پیش‌بینی «سری زمانی مکی گلاس» پیداست، روش «گرادیان نزولی» در حذف اثر نویز با دامنه بالا<sup>۱</sup> در بیشتر موارد نسبت به روش «الگوریتم خوشه‌بندی K-Means» بهتر عمل می‌کند. البته این بهتر بودن به قیمت حجم بالای محاسباتی که منجر به افزایش زمان اجرای برنامه می‌گردد به دست آمده است.

زمان اجرای برنامه، وابسته به الف) تعداد پارامترهای انتخاب شده برای آموزش ب) حجم محاسبات ج) برخی ویژگی‌های روش آموزش مانند پیچیدگی الگوریتم آن؛ می‌باشد. به عنوان مثال با افزایش تعداد نرون‌های لایه پنهان؛ تعداد پارامترهای انتخاب شده برای آموزش و در نتیجه زمان اجرای برنامه بیشتر می‌گردد. هم‌چنین محاسبات بازه‌ای در لایه پنهان و لایه خروجی به شدت بر حجم محاسبات می‌افزاید. این قبیل محاسبات در روش «گرادیان نزولی» حضور پر رنگتری دارند. و یا به دلیل کمتر بودن پیچیدگی «الگوریتم خوشه‌بندی K-Means» نسبت به «گرادیان نزولی»؛ آموزش شبکه به کمک این الگوریتم؛ زمان اجرای کمتری دارد. در جدول ۷، نسبت زمان اجرای برنامه با روش «الگوریتم خوشه‌بندی K-Means» به زمان اجرای برنامه با روش «گرادیان نزولی» برای تعداد متفاوتی از نرون‌های لایه پنهان آورده شده است.

جدول ۷. نسبت زمان اجرای برنامه با روش «الگوریتم خوشه‌بندی K-Means» به زمان اجرای برنامه با روش «گرادیان نزولی» (زمان‌ها بر حسب ثانیه است)

| تعداد نرون‌های لایه پنهان | سیستم غیر خطی پویای $U$ شکل با پنج ورودی | سری زمانی مکی گلاس         |
|---------------------------|------------------------------------------|----------------------------|
| ۲۰۰                       | $\frac{41.6}{120.1} = 0.35$              | $\frac{30}{73.4} = 0.41$   |
| ۲۰                        | $\frac{14.2}{37.5} = 0.38$               | $\frac{11.5}{26.7} = 0.43$ |
| ۱۰                        | $\frac{13.1}{33.4} = 0.39$               | $\frac{10}{22.5} = 0.44$   |

## ۶- نتیجه‌گیری

همان‌طور که از نتایج معلوم می‌گردد صرف‌نظر از روش آموزش؛ عملکرد شبکه عصبی RBF گرانولی نسبت به شبکه عصبی RBF، در شرایطی که داده‌های ورودی آغشته به نویز هستند بهتر می‌باشد. هم‌چنین مشاهده می‌گردد که هر چه میزان نویز تزریق شده به داده‌ها بیشتر باشد درصد بهبود عملکرد شبکه عصبی RBF گرانولی نسبت به شبکه عصبی RBF افزایش می‌یابد. به عبارت دیگر بازه‌ای بودن پارامترها سیستم را در

مقابل نویز، قوی<sup>۳</sup> می‌گرداند. این مسأله بدان سبب است که تابع فعال‌ساز گرانولی در ساختار شبکه عصبی RBF، با تنظیم «پارامترهای مربوط به پهنای عدم قطعیت<sup>۴</sup>» یعنی کران پائین و کران بالای انحراف استاندارد، باعث کاهش حساسیت به تغییرات ورودی می‌شود. بنابراین توصیه می‌گردد که اگر نویز زیادی در داده‌ها وجود دارد با توجه به بهتر بودن عملکرد شبکه عصبی RBF گرانولی نسبت به شبکه عصبی RBF، از شبکه عصبی RBF گرانولی استفاده شود. اما اگر نویز موجود؛ مقدار قابل توجهی نداشته باشد، با توجه به سادگی ساختار شبکه عصبی RBF از این ساختار استفاده گردد. سادگی ساختار شبکه عصبی RBF که به دلیل رهایی از محاسبات بازه‌ای در لایه‌های میانی و خروجی می‌باشد، سبب کم شدن خطای ناشی از این محاسبات و هم‌چنین کاهش حجم محاسباتی برنامه می‌گردد. کم شدن مقدار تقریب در محاسبه خروجی، می‌تواند در شرایط بدون نویز و به ویژه در شناسایی سیستم‌های غیر خطی پویا؛ منجر به اندکی بهبود در نتایج گردد.

از طرفی نتایج به دست آمده از شبکه عصبی RBF گرانولی، در شرایطی که داده‌های ورودی آغشته به نویز قوی نیستند؛ نشان دهنده این است که عملکرد «الگوریتم خوشه‌بندی K-Means» در شناسایی سیستم‌های غیر خطی پویا که آشوبی نیستند بهتر است. البته برای شناسایی و پیش‌بینی توابع آشوب برای افق زمانی بزرگ، نتایج به دست آمده بیانگر آن است که عملکرد شبکه عصبی RBF گرانولی با آموزش پارامترهای لایه میانی بوسیله «الگوریتم گرادیان نزولی» عملکرد بهتری دارد. هر چند که «الگوریتم خوشه‌بندی K-Means» نیز در این شرایط با کیفیت پایین‌تری کاربرد دارد.

با افزایش قدرت نویز در داده‌های ورودی برتری «الگوریتم گرادیان نزولی» به تدریج افزایش می‌یابد تا جایی که در نویز با «نسبت سیگنال به نویز» برابر صفر، عملکرد شبکه عصبی RBF گرانولی با آموزش پارامترهای لایه میانی بوسیله روش «گرادیان نزولی»؛ برای شناسایی سیستم‌های غیر خطی پویا و شناسایی و پیش‌بینی توابع آشوب، بهتر از حالتی است که پارامترهای لایه میانی بوسیله روش «الگوریتم خوشه‌بندی K-Means» آموزش می‌بینند. چرا که در هنگام وجود نویز قوی در داده ورودی، تعیین تعلق آن به یک خوشه دشوار می‌باشد و خطای بوجود آمده در تعیین خوشه بر تقریب ناشی از محاسبات بازه‌ای غلبه می‌یابد.

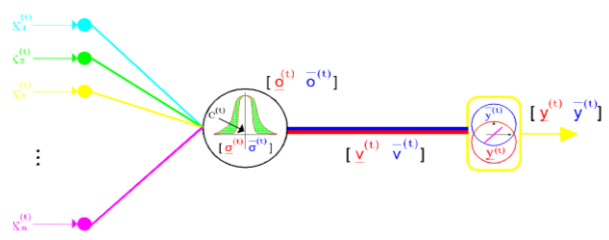
با این وجود، زمان اجرای برنامه با آموزش پارامترهای لایه میانی بوسیله «الگوریتم خوشه‌بندی K-Means»، کمتر از روش گرادیان نزولی است. علت آن این است که مرتبه پیچیدگی الگوریتم و حجم محاسباتی برنامه کاهش یافته و همگرایی پارامترها سریع‌تر صورت می‌پذیرد. از اینرو در شرایطی که داده‌های ورودی آغشته به نویز قوی نیستند؛ در تعداد مرحله مساوی، نسبت به روش «گرادیان نزولی»؛ به جواب بهتری خواهیم رسید.

<sup>3</sup> Robust<sup>4</sup> The parameters responsible for the width of uncertainly<sup>1</sup> Run time<sup>2</sup> High Amplitude

## ۷- پیوست: اثبات بهبود یافتن عملکرد شبکه عصبی RBF گرانولی نسبت به شبکه عصبی RBF در شرایط نویزی

۱-۷ اثبات برای حالتی که یک نرون در لایه میانی وجود داشته باشد

برای اثبات بهبود یافتن عملکرد شبکه عصبی RBF گرانولی نسبت به شبکه عصبی RBF در شرایط نویزی؛ از یک شبکه عصبی RBF گرانولی با یک نرون در لایه میانی استفاده می‌کنیم. اثبات می‌گردد که هر چه میزان نویز تزریق شده به داده‌ها بیشتر باشد، شبکه عصبی RBF گرانولی با افزایش «فاصله کران پائین انحراف استاندارد از کران بالای آن» و در نتیجه افزایش «پهنای عدم قطعیت و یا همان تابع فعال‌ساز گاوسی گرانولی»، اثر نویز در خروجی شبکه را مهار می‌نماید. هم‌چنین هر چه میزان نویز تزریق شده به داده‌ها کمتر باشد ساختار شبکه عصبی RBF گرانولی به شبکه عصبی RBF نزدیکتر می‌شود تا در شرایط بدون نویز دو ساختار مشابه یکدیگر می‌گردند. البته این تغییر ساختار ناشی از بازه‌ای بودن دو پارامتر سیستم، یکی «انحراف استاندارد تابع فعال‌ساز گرانولی» در لایه میانی؛ و دیگری «وزن‌های بازه‌ای» در لایه خروجی؛ می‌باشد.



شکل ۹. یک شبکه عصبی RBF گرانولی با یک نرون در لایه میانی

محاسبه کران پائین خروجی یک تابع گاوسی گرانولی در لایه میانی به شکل زیر است:

$$\underline{o}^{(t)} = \psi \left[ \sum_{r=1}^n \left( \frac{x_r^{(t)} + n_r^{(t)} - c_r^{(t)}}{\underline{\sigma}_r^{(t)}} \right)^2 \right] = \psi \left[ \sum_{r=1}^n \left( \bar{u}^{(t)} \right)^2 \right] = \psi \left( \bar{d}^{(t)} \right) = \exp \left( -\frac{1}{2} \bar{d}^{(t)} \right) \quad (79)$$

اکنون برای محاسبه کران بالای خروجی یک تابع گاوسی گرانولی در لایه میانی خواهیم داشت:

$$\bar{o}^{(t)} = \psi \left[ \sum_{r=1}^n \left( \frac{x_r^{(t)} + n_r^{(t)} - c_r^{(t)}}{\bar{\sigma}_r^{(t)}} \right)^2 \right] = \psi \left[ \sum_{r=1}^n \left( \underline{u}^{(t)} \right)^2 \right] = \psi \left( \underline{d}^{(t)} \right) = \exp \left( -\frac{1}{2} \underline{d}^{(t)} \right) \quad (80)$$

همان‌طور که در محاسبه خروجی نهایی شبکه عصبی RBF گرانولی دیدیم؛ برای محاسبه خروجی نهایی این شبکه با یک نرون هم چهار حالت زیر را خواهیم داشت.

- حالت اول: اگر  $v^{(t)} > 0$  و  $v^{(t)} > |s^{(t)}|$  باشد آن‌گاه  $\underline{v}^{(t)}, \bar{v}^{(t)} > 0$  خواهد بود.
- حالت دوم: اگر  $v^{(t)} > 0$  و  $v^{(t)} < |s^{(t)}|$  باشد آن‌گاه  $\underline{v}^{(t)} < 0$  و  $\bar{v}^{(t)} > 0$  خواهد بود.
- حالت سوم: اگر  $v^{(t)} < 0$  و  $|v^{(t)}| > |s^{(t)}|$  باشد آن‌گاه  $\underline{v}^{(t)}, \bar{v}^{(t)} < 0$  خواهد بود.
- حالت چهارم: اگر  $v^{(t)} < 0$  و  $|v^{(t)}| < |s^{(t)}|$  باشد آن‌گاه  $\underline{v}^{(t)} < 0$  و  $\bar{v}^{(t)} > 0$  خواهد بود.

در حالت‌های اول، دوم و چهارم، کران پایین و کران بالای خروجی نهایی شبکه عصبی RBF گرانولی با یک نرون، به همراه پارامترهای خطی کران پایین و کران بالای لایه خروجی شبکه به این صورت خواهد بود:

$$\underline{y_N}^{(t)} = \underline{v}^{(t)} \times \underline{o}^{(t)} = (v^{(t)} - |s^{(t)}|) \times \underline{o}^{(t)} \quad (81)$$

$$\bar{y_N}^{(t)} = \bar{v}^{(t)} \times \bar{o}^{(t)} = (v^{(t)} + |s^{(t)}|) \times \bar{o}^{(t)} \quad (82)$$

در حالت سوم به علت منفی بودن  $\underline{v}^{(t)}$  و مثبت بودن  $\bar{v}^{(t)}$  و  $\bar{o}^{(t)}$ ، محاسبه کران پایین و کران بالای خروجی نهایی شبکه عصبی RBF گرانولی با یک نرون، به همراه پارامترهای خطی کران پایین و کران بالای لایه خروجی شبکه به این صورت خواهد بود:

$$\underline{y_N}^{(t)} = \underline{v}^{(t)} \times \underline{o}^{(t)} = (v^{(t)} - |s^{(t)}|) \times \bar{o}^{(t)} \quad (83)$$

$$\bar{y_N}^{(t)} = \bar{v}^{(t)} \times \bar{o}^{(t)} = (v^{(t)} + |s^{(t)}|) \times \underline{o}^{(t)} \quad (84)$$

و خروجی نهایی از میانگین کران پائین و کران بالا به دست می‌آید:

$$y_N^{(t)} = \frac{\underline{y_N}^{(t)} + \bar{y_N}^{(t)}}{2} \quad (85)$$

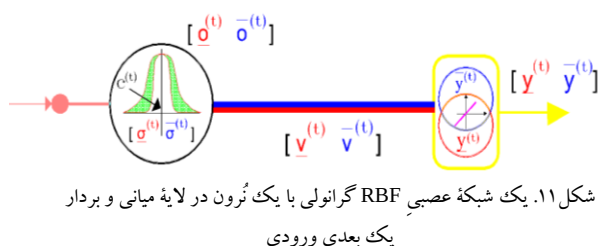
اکنون با سه فرض زیر اثبات را ارائه می‌کنیم:

- وزن لایه خروجی، یعنی  $v^{(t)}$  بازه‌ای نبوده و مساوی یک باشد، در آن صورت خواهیم داشت:

$$y_N^{(t)} = \frac{\underline{y_N}^{(t)} + \bar{y_N}^{(t)}}{2} = \frac{(\bar{v}^{(t)} \times \bar{o}^{(t)}) + (\underline{v}^{(t)} \times \underline{o}^{(t)})}{2} = \frac{\bar{o}^{(t)} + \underline{o}^{(t)}}{2} = \frac{\psi \left[ \sum_{r=1}^n \left( \frac{x_r^{(t)} + n_r^{(t)} - c_r^{(t)}}{\underline{\sigma}_r^{(t)}} \right)^2 \right] + \psi \left[ \sum_{r=1}^n \left( \frac{x_r^{(t)} + n_r^{(t)} - c_r^{(t)}}{\bar{\sigma}_r^{(t)}} \right)^2 \right]}{2} \quad (86)$$

- هم‌چنین مقدار اولیه‌ی مرکز تابع فعال‌ساز گاوسی نرون لایه پنهان یعنی  $C$  مساوی صفر در نظر گرفته می‌شود.

اکنون برای کاهش تعداد متغیرها در تابع خروجی در وضعیت بدون نویز و نویزی، با فرض این که ابعاد بردار ورودی برابر یک باشد یعنی  $n = 1$  اثبات را ادامه می‌دهیم:



در این صورت خواهیم داشت:

$$y^{(t)} = f(x^{(t)}) = \frac{1}{2} \times \psi \left[ (x^{(t)})^2 \right] + \frac{1}{2} \times \psi \left[ \frac{(x^{(t)})^2}{(1+\sigma^{(t)})^2} \right] \quad (91)$$

$$y_{\mathcal{N}}^{(t)} = f(x^{(t)}, n^{(t)}) = \frac{1}{2} \times \psi \left[ (x^{(t)} + n^{(t)})^2 \right] + \frac{1}{2} \times \psi \left[ \frac{(x^{(t)} + n^{(t)})^2}{(1+\sigma^{(t)})^2} \right] \quad (92)$$

تفاضل این دو تابع، تابعی است که اختلال ایجاد شده بوسیله نویز<sup>۱</sup> را نشان می‌دهد و عبارت است از:

$$\begin{aligned} DCN &= |y_{\mathcal{N}}^{(t)} - y^{(t)}| \\ &= \frac{1}{2} \\ &\times \left| \psi \left[ (x^{(t)} + n^{(t)})^2 \right] + \psi \left[ \frac{(x^{(t)} + n^{(t)})^2}{(1+\sigma^{(t)})^2} \right] - \psi \left[ (x^{(t)})^2 \right] - \psi \left[ \frac{(x^{(t)})^2}{(1+\sigma^{(t)})^2} \right] \right| \end{aligned} \quad (93)$$

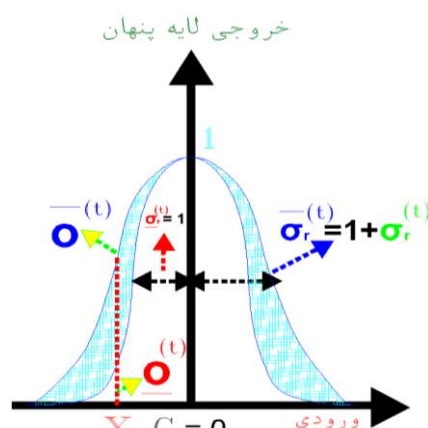
$$\begin{aligned} DCN &= \frac{1}{2} \times \left| \exp \left[ \frac{-(x^{(t)} + n^{(t)})^2}{2} \right] + \exp \left[ \frac{-(x^{(t)} + n^{(t)})^2}{2(1+\sigma^{(t)})^2} \right] - \exp \left[ \frac{-(x^{(t)})^2}{2} \right] - \exp \left[ \frac{-(x^{(t)})^2}{2(1+\sigma^{(t)})^2} \right] \right| \end{aligned} \quad (94)$$

همان‌طور که در بخش ۲ گفته شد؛ تابع توزیع احتمال گاوسی یا نرمال<sup>۲</sup> بر اساس قضیه حد مرکزی از اهمیت زیادی برخوردار است. اکنون «تابع توزیع احتمال» برای نویز ورودی را که به آن «تابع توزیع

$$y_{\mathcal{N}}^{(t)} = \frac{\psi \left[ \sum_{r=1}^n \left( \frac{x_r^{(t)} + n_r^{(t)}}{\sigma_r^{(t)}} \right)^2 \right] + \psi \left[ \sum_{r=1}^n \left( \frac{x_r^{(t)} + n_r^{(t)}}{1 + \sigma_r^{(t)}} \right)^2 \right]}{2} \quad (87)$$

• کران پائین انحراف استاندارد مساوی یک ( $\sigma_r^{(t)} = 1$ ) و کران بالای انحراف استاندارد به صورت  $\bar{\sigma}_r^{(t)} = 1 + \sigma_r^{(t)}$  در نظر گرفته می‌شود که  $\sigma_r^{(t)} = [0 \quad 3]$  بازه تغییرات انحراف استاندارد با گام ۰/۰۱ خواهد بود.

$$\begin{aligned} y_{\mathcal{N}}^{(t)} &= \frac{\psi \left[ \sum_{r=1}^n \left( \frac{x_r^{(t)} + n_r^{(t)}}{1} \right)^2 \right] + \psi \left[ \sum_{r=1}^n \left( \frac{x_r^{(t)} + n_r^{(t)}}{1 + \sigma_r^{(t)}} \right)^2 \right]}{2} = \frac{1}{2} \times \\ &\psi \left[ \sum_{r=1}^n (x_r^{(t)} + n_r^{(t)})^2 \right] + \frac{1}{2} \times \psi \left[ \sum_{r=1}^n \frac{(x_r^{(t)} + n_r^{(t)})^2}{(1 + \sigma_r^{(t)})^2} \right] \end{aligned} \quad (88)$$



شکل ۱۰. تابع فعال‌ساز گاوسی گرانولی در نرون لایه میانی و پیش از آغاز آموزش مرکز تابع فعال‌ساز

سپس خروجی شبکه بدون افزودن نویز ورودی و با فرضیات بالا و در نظر گرفتن بردار  $X$  به عنوان متغیر مستقل به شکل زیر به دست می‌آید:

$$\begin{aligned} y^{(t)} &= f_{(t)}(X) = f(x_1^{(t)}, x_2^{(t)}, x_3^{(t)}, \dots, x_n^{(t)}) \\ &= \frac{1}{2} \times \psi \left[ \sum_{r=1}^n (x_r^{(t)})^2 \right] + \frac{1}{2} \\ &\times \psi \left[ \sum_{r=1}^n \frac{(x_r^{(t)})^2}{(1 + \sigma_r^{(t)})^2} \right] \end{aligned} \quad (89)$$

با افزودن نویز به داده‌ها، معادله خروجی به صورت زیر خواهد بود:

$$\begin{aligned} y_{\mathcal{N}}^{(t)} &= f_{(t)}(X, \mathcal{N}) \\ &= f(x_1^{(t)}, x_2^{(t)}, x_3^{(t)}, \dots, x_n^{(t)}, n_1^{(t)}, n_2^{(t)}, n_3^{(t)}, \dots, n_n^{(t)}) \\ &= \frac{1}{2} \times \psi \left[ \sum_{r=1}^n (x_r^{(t)} + n_r^{(t)})^2 \right] + \frac{1}{2} \\ &\times \psi \left[ \sum_{r=1}^n \frac{(x_r^{(t)} + n_r^{(t)})^2}{(1 + \sigma_r^{(t)})^2} \right] \end{aligned} \quad (90)$$

<sup>1</sup> Distortion Caused by Noise (DCN)

<sup>2</sup> Gaussian Probability Distribution Function (Gaussian PDF)

$$DCN = |y_N - y| = \frac{1}{2} \times \left| \exp \left[ \frac{-(x+n)^2}{2} \right] + \exp \left[ \frac{-(x+n)^2}{2(1+\sigma)^2} \right] - \exp \left[ \frac{-x^2}{2} \right] - \exp \left[ \frac{-x^2}{2(1+\sigma)^2} \right] \right| \quad (99)$$

متغیر تصادفی  $n$  روی فضای احتمال<sup>۴</sup>  $(\Omega, \mathcal{F}, P)$  تعریف می‌گردد که در آن داریم [۴۱]:

- فضای نمونه<sup>۵</sup> شامل مقادیر تصادفی نویز و در واقع همان نتایج به دست آمده از آزمایشات است.
- فضای رویداد<sup>۶</sup> و یا فضای فرآیند تصادفی<sup>۷</sup> است که بخشی از فضای نمونه را در بر می‌گیرد و شامل رویدادهایی است که وقوع پذیرند و در واقع احتمال وقوع بزرگتر از صفر دارند. به عنوان مثال رویداد قرار گرفتن «بزرگی نویز در بازه  $[-0.6745\sigma_n, 0.6745\sigma_n]$ » احتمال وقوع بالایی دارد.
- تابع اندازه احتمال رویداد<sup>۸</sup> است که وظیفه آن اختصاص یک عدد حقیقی به رویداد می‌باشد و در واقع فضای رویداد را به فضای اعداد حقیقی نگاشت می‌کند. می‌توان از تابع چگالی احتمال به عنوان تابع اندازه احتمال رویداد استفاده نمود.

مقدار مورد انتظار یا امید ریاضی برای اختلال ایجاد شده بوسیله نویز؛ عبارت است از:

$$E(\mathcal{N}) = \int_{\Omega} DCN \times dP = \int_{-\infty}^{\infty} DCN \times f_N(n) dn = \int_{-\infty}^{\infty} |y_N - y| \times f_N(n) dn \quad (100)$$

با تقریب خوبی می‌توان حدود انتگرال‌گیری را به شکل زیر عوض نمود:

$$E(\mathcal{N}) = \int_{-4\sigma}^{4\sigma} DCN \times f_N(n) dn = \int_{-4\sigma}^{4\sigma} |y_N - y| \times f_N(n) dn \quad (101)$$

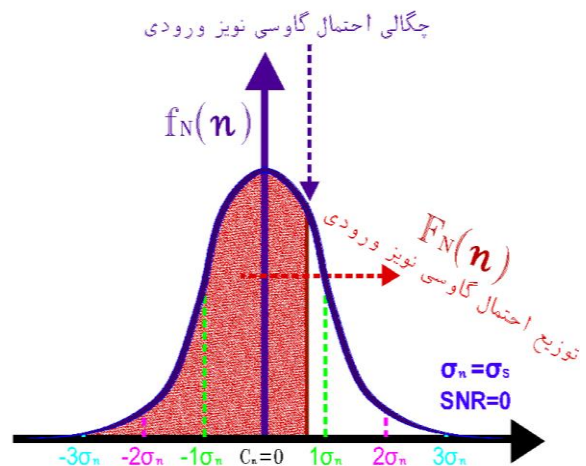
می‌توان فرم پیوسته تابع امید ریاضی یعنی  $E(\mathcal{N})$  را به صورت زیر گسسته نمود. در حالت گسسته این تابع از طریق جمع مقادیر نمونه‌ها و

تجمعی<sup>۱</sup> نیز گفته می‌شود، از نوع «نرمال یا گاوسی با میانگین صفر» و به صورت  $\mathcal{N}(0, \sigma_n^2)$  در نظر می‌گیریم و به شکل زیر تعریف می‌نمائیم [۳۹-۴۱]:

$$F_N(n) = \int_{-\infty}^n f_N(z) dz = \frac{1}{\sqrt{2\pi} \times \sigma_n} \int_{-\infty}^n \psi \left( \frac{z^2}{\sigma_n^2} \right) dz = \frac{1}{\sqrt{2\pi} \times \sigma_n} \int_{-\infty}^n \exp \left( -\frac{1}{2} \times \frac{z^2}{\sigma_n^2} \right) dz \quad (95)$$

تابع چگالی احتمال<sup>۲</sup> برای نویز ورودی به شکل زیر به دست می‌آید.

$$f_N(n) = \frac{d}{dn} F_N(n) = \frac{1}{\sqrt{2\pi} \times \sigma_n} \psi \left( \frac{n^2}{\sigma_n^2} \right) = \frac{1}{\sqrt{2\pi} \times \sigma_n} \exp \left( -\frac{1}{2} \times \frac{n^2}{\sigma_n^2} \right) \quad (96)$$



شکل ۱۲. تابع توزیع احتمال و تابع چگالی احتمال گاوسی برای نویز ورودی

همان طور که در شکل ۹ می‌بینیم، فرض بر این است که قدرت نویز زیاد بوده و با قدرت سیگنال برابر باشد یعنی خواهیم داشت  $SNR = 0$  (dB) و در نتیجه [۱۱]:

$$SNR = 10 \log_{10} \frac{\text{the variance of the signal}}{\text{the variance of the noise}} = 10 \log_{10} \frac{\sigma_s^2}{\sigma_n^2} = 0 \rightarrow \sigma_s = \sigma_n \quad (97)$$

آنجا که  $\sigma_s^2$  واریانس سیگنال و  $\sigma_n^2$  واریانس نویز می‌باشد. اکنون فرض می‌کنیم  $\sigma_s = \sigma_n = \sigma' = 0.33$  باشد، در آن صورت تابع چگالی احتمال برای نویز که در حالت گسسته به آن تابع جرم احتمال<sup>۳</sup> گویند به شکل زیر ساده می‌گردد:

$$f_N(n) = \frac{1}{\sqrt{2\pi} \times \sigma'} \exp \left( -\frac{1}{2} \times \frac{n^2}{\sigma'^2} \right) = \frac{1}{2.5 \times 0.33} e^{-\frac{n^2}{2 \times 0.33^2}} = \frac{1}{0.825} e^{-\frac{n^2}{0.2178}} = 1.2 e^{-4.6 n^2} \quad (98)$$

اکنون با فرض داشتن یک نمونه ورودی یعنی  $T=1$ ، می‌توانیم تابع دو متغیره  $DCN$  را به یک تابع یک متغیره با متغیر مستقل  $n$  تبدیل نمائیم.

<sup>4</sup> Probability space

<sup>5</sup> Sample space

<sup>6</sup> Event space

<sup>7</sup> Stochastic Process space

<sup>8</sup> Event Probability Function

<sup>1</sup> Cumulative Distribution Function (CDF)

<sup>2</sup> Probability Density Function (pdf)

<sup>3</sup> Probability Mass Function (pmf)



۲-۷ تعمیم اثبات فوق برای حالتی که  $m$  نرون در لایه میانی وجود داشته باشد

**تعریف ۱.** فرآیند تصادفی شمارا<sup>۲</sup> (با زمان گسسته<sup>۳</sup>) روی فضای احتمال  $(\Omega, \mathcal{F}, P)$  عبارت است از دنباله‌ای از متغیرهای تصادفی مانند  $\{n_r^{(t)}\}_{r \in \mathbb{N}}$  روی این فضا و آن را به صورت  $\{n_r^{(t)}\}_{r \in \mathbb{N}}$  نشان می‌دهیم [۳۹].

**قرارداد.** از این پس منظور از  $\mathbb{R}^\infty$  یعنی مجموعه همه دنباله‌های اعداد حقیقی؛ و یا به عبارت دیگر مجموعه همه توابع از  $\mathbb{N}$  به  $\mathbb{R}$  که هر تابع می‌تواند یک دنباله را تولید کند. به عنوان مثال تابع نویز  $n_r^{(t)} = f(r, t)$  با «تابع توزیع احتمال گاوسی» می‌تواند یکی از این توابع باشد [۴۰-۳۸].

**تعریف ۲.** مجموعه  $A$  در  $\mathbb{R}^\infty$  را یک بازه با بُعد متناهی گوئیم اگر  $n \in N$  و بازه‌های متناهی و نامتناهی  $I_1, I_2, I_3, \dots, I_n$  موجود باشند که  $A = \{N^{(t)} \in \mathbb{R}^\infty : n_r^{(t)} \in I_r, r = 1, 2, \dots, n\}$  به عنوان مثال مقدار نویز افزوده شده به ورودی  $r$  ام، یعنی  $x_r^{(t)}$  که به صورت  $n_r^{(t)}$  نشان داده می‌شود، در بازه  $[-4\sigma_n, 4\sigma_n]$  که یک بازه با بُعد متناهی است قرار دارد. بُعد متناهی این بازه برابر بُعد بردار نویز ورودی یعنی  $N^{(t)} = [n_1^{(t)} \ n_2^{(t)} \ \dots \ n_n^{(t)}]$  می‌باشد. رده تمام بازه-های با بُعد متناهی را با  $G$  نشان می‌دهند [۲۷].

**تعریف ۳.** فرض می‌کنیم که مجموعه  $B$  به صورت زیر تعریف گردد [۴۰-۳۸]:

$$B = \{X^{(t)}, N^{(t)}\} \\ \in \mathbb{R}^\infty: (x_1^{(t)}, x_2^{(t)}, x_3^{(t)}, \dots, x_n^{(t)}, n_1^{(t)}, n_2^{(t)}, n_3^{(t)}, \dots, n_n^{(t)}) \in B_n\} \quad (106)$$

- مجموعه  $B$  در  $\mathbb{R}^\infty$  را یک استوانه<sup>۴</sup> با بُعد متناهی گوئیم.
- رده تمام استوانه‌های با بُعد متناهی با  $e_\infty$  نشان داده می‌شود که  $e_\infty$  یک میدان است.
- در این صورت  $e_\infty$  می‌تواند یک میدان سیگمایی<sup>۵</sup> باشد.
- از دیدگاه احتمال، میدان سیگمایی را می‌توان به عنوان اطلاعات داده شده یک آزمایش در نظر گرفت.
- مجموعه‌های بُر روی  $\Omega$  و  $P$  از فضای احتمال  $(\Omega, \mathcal{F}, P)$  به شکل  $\sigma(e_\infty)$  و  $\sigma(G_\infty)$  تعریف می‌گردند.

هنجارسازی<sup>۱</sup> حاصل جمع نسبت به تعداد نمونه‌ها محاسبه می‌شود و یا می‌توان از تابع چگالی احتمال نویز ورودی به صورت زیر استفاده نمود.

$$E(N) = \sum_{n=-4\sigma}^{4\sigma} (|y_N - y| \times f_N(n) \Delta n) \quad (102)$$

هر چه مقدار  $\Delta n$  را کوچکتر در نظر بگیریم دقت محاسبه  $E(N)$  بیشتر خواهد بود.  $\Delta n$  را می‌توان برابر  $0.001$  یا  $0.0001$  در نظر گرفت. اکنون تابع  $E(N^2)$  را به شکل زیر به دست می‌آوریم:

$$E(N^2) = \int_{-4\sigma}^{4\sigma} (DCN)^2 \times f_N(n) dn = \int_{-4\sigma}^{4\sigma} (y_N - y)^2 \times f_N(n) dn \quad (103)$$

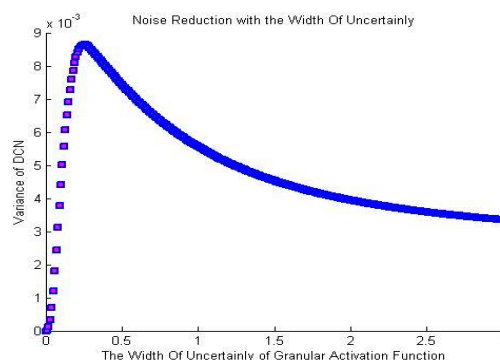
می‌توان فرم پیوسته تابع  $E(N^2)$  را به صورت زیر گسسته نمود:

$$E(N^2) = \sum_{n=-4\sigma}^{4\sigma} [(y_N - y)^2 \times f_N(n) \Delta n] \quad (104)$$

واریانس تابع  $DCN$ ، به شکل زیر محاسبه می‌گردد:

$$Var(N) = E(N^2) - [E(N)]^2 \quad (105)$$

نمودار تغییرات واریانس تابع  $DCN$  بر حسب  $\sigma$  مانند شکل ۱۰ به دست می‌آید. این شکل نشان می‌دهد که ابتدا با افزایش مقدار  $\sigma$  تغییر انحراف استاندارد تابع فعال‌ساز گرانولی تا سطح  $\sigma' = 0.33$  (انحراف استاندارد نویز)، واریانس تابع  $DCN$  افزایش می‌یابد. چرا که پهنای عدم قطعیت در تابع فعال‌ساز گرانولی هنوز قدرت خنثی نمودن اثر نویز را ندارد و تغییرات مقدار نویز سبب تغییر اختلال ایجاد شده بوسیله نویز می‌گردد. که آن نیز به نوبه خود تغییر در واریانس این اختلال را در پی دارد که در این مرحله افزایشی می‌باشد. با افزایش بیشتر مقدار  $\sigma$ ، پهنای عدم قطعیت در تابع فعال‌ساز گرانولی بر نویز ورودی غلبه نموده و واریانس تابع  $DCN$  رو به کاهش می‌گذارد. کاهش واریانس تابع  $DCN$  به معنی آن است که با افزایش پهنای عدم قطعیت در تابع فعال‌ساز گرانولی، مقدار تابع  $DCN$  حول میانگین آن پراکندگی کمتری دارد و چون تابع توزیع احتمال نویز ورودی، گاوسی یا نرمال با میانگین صفر می‌باشد، پس تابع  $DCN$  بیشتر به سمت صفر میل کرده و در واقع اختلال ایجاد شده بوسیله نویز کاهش یافته است.



شکل ۱۳. تغییرات واریانس تابع  $DCN$  بر حسب پهنای عدم قطعیت یعنی  $\sigma$

<sup>2</sup> Countable Stochastic Process

<sup>3</sup> Discrete Time

<sup>4</sup> Cylinder

<sup>5</sup>  $\sigma$ -field

<sup>1</sup> Normalizing

- مجموعه  $B_\infty$  را به صورت  $B_\infty = \sigma(G_\infty) = \sigma(e_\infty)$  تعریف می‌کنیم.

**گزاره:** اثبات می‌گردد که هر فرآیند تصادفی مانند فرآیند تصادفی شمارای  $\{n_r^{(t)}\}_{r \in N}$  یک بردار تصادفی به صورت  $(\Omega, \mathcal{F}, P) \rightarrow (\mathbb{R}^\infty, B_\infty)$  به دست می‌دهد [۳۹-۴۱].

**قضیه توسعه کلموگورف<sup>۱</sup>:** فرض کنید که برای هر  $n \in N$  یک اندازه احتمال  $P_n$  روی  $(\mathbb{R}^\infty, B)$  در دست است. می‌خواهیم یک اندازه احتمال  $P$  روی  $(\mathbb{R}^\infty, B_\infty)$  طوری بسازیم که تحدید  $P$  به  $B_n$  همان  $P_n$  شود. یعنی داشته باشیم [۳۹-۴۱]:

$$\forall B \in B_n; P\left(\left\{X^{(t)}, \mathcal{N}^{(t)}\right\} \in \mathbb{R}^\infty: (x_1^{(t)}, x_2^{(t)}, x_3^{(t)}, \dots, x_n^{(t)}, n_1^{(t)}, n_2^{(t)}, n_3^{(t)}, \dots, n_n^{(t)}) \in B\right) = P_n(B) \quad (107)$$

می‌دانیم که  $\mathcal{F} \subset \Omega$  است. اگر مجموعه  $F$  را به عنوان مجموعه‌ای از درایه‌های بردار  $\mathcal{N}^{(t)}$  در نظر بگیریم و  $F$  یک مجموعه شمارا و زیر مجموعه  $R$  باشد. و با توجه به این که می‌دانیم؛ دنباله  $(n_r^{(t)})_{r \in N}$  دنباله‌ای از متغیرهای تصادفی است آن‌گاه در صورت برقراری شرایط زیر:

- $P(\cdot) \geq 0$ , (108)
- $\sum_{n_{n+1}^{(t)} \in F} P(n_1^{(t)}, n_2^{(t)}, n_3^{(t)}, \dots, n_{n+1}^{(t)}) = P(n_1^{(t)}, n_2^{(t)}, n_3^{(t)}, \dots, n_n^{(t)})$ , (109)
- $\sum_{n_1^{(t)} \in F} P(n_1^{(t)}) = 1$ , (110)

فرآیند شمارای  $\{x_r^{(t)}\}_{r \in N}$  که می‌تواند دارای ویژگی غیرخطی و یا آشوبی باشد؛ چنان موجود است که:

$$P(x_1^{(t)} \leftarrow x_1^{(t)} + n_1^{(t)}, \dots, x_n^{(t)} \leftarrow x_n^{(t)} + n_n^{(t)}) = P(n_1^{(t)}, n_2^{(t)}, n_3^{(t)}, \dots, n_n^{(t)}) \quad (111)$$

برای تعمیم تعداد ورودی‌ها از یک ورودی به چند ورودی؛ از قضیه توسعه کلموگورف و اثبات به کمک استقراء ریاضی؛ به شکل زیر استفاده می‌کنیم:

- اگر برای  $1 \in N$  یک اندازه احتمال  $P_1$  روی  $(\mathbb{R}^\infty, B)$  در دست داشته باشیم؛ می‌توانیم یک اندازه احتمال  $P$  روی  $(\mathbb{R}^\infty, B_\infty)$  طوری بسازیم که تحدید  $P$  به  $B_1$  همان  $P_1$  شود. یعنی داشته باشیم:
- $$\forall B \in B_1; P\left(\left\{X^{(t)}, \mathcal{N}^{(t)}\right\} \in \mathbb{R}^\infty: (x_1^{(t)}, n_1^{(t)}) \in B\right) = P_1(B) \quad (112)$$

- اگر برای  $r \in N$  یک اندازه احتمال  $P_r$  روی  $(\mathbb{R}^\infty, B)$  در دست داشته باشیم؛ می‌توانیم یک اندازه احتمال  $P$  روی  $(\mathbb{R}^\infty, B_\infty)$  طوری بسازیم که تحدید  $P$  به  $B_r$  همان  $P_r$  شود. یعنی داشته باشیم:
- $$\forall B \in B_r; P\left(\left\{X^{(t)}, \mathcal{N}^{(t)}\right\} \in \mathbb{R}^\infty: (x_1^{(t)}, x_2^{(t)}, n_1^{(t)}, n_2^{(t)}) \in B\right) = P_r(B) \quad (113)$$

- اگر برای  $r+1 \in N$  یک اندازه احتمال  $P_{r+1}$  روی  $(\mathbb{R}^\infty, B)$  در دست داشته باشیم؛ می‌توانیم یک اندازه احتمال  $P$  روی  $(\mathbb{R}^\infty, B_\infty)$  طوری بسازیم که تحدید  $P$  به  $B_{r+1}$  همان  $P_{r+1}$  شود. یعنی داشته باشیم:

$$\forall B \in B_{r+1}; P\left(\left\{X^{(t)}, \mathcal{N}^{(t)}\right\} \in \mathbb{R}^\infty: (x_1^{(t)}, x_2^{(t)}, n_1^{(t)}, n_2^{(t)}) \in B\right) = P_{r+1}(B) \quad (114)$$

پس با هر تعداد ورودی قادر به ساختن یک اندازه احتمال  $P$  روی  $(\mathbb{R}^\infty, B_\infty)$  خواهیم بود.

**قضیه وجود نگاشت<sup>۲</sup> شبکه عصبی پیشرو<sup>۳</sup> کلموگورف:** هر تابع پیوسته  $f: [0 \ 1]^n \rightarrow \mathbb{R}^m$  را می‌توان به طور دقیق توسط یک شبکه عصبی سه لایه‌ای پیشرو<sup>۴</sup>، که دارای  $n$  نرون در لایه ورودی و  $(2n+1)$  نرون در لایه پنهان و  $m$  نرون در لایه خروجی می‌باشد تحقق بخشید. طبق این قضیه چنین شبکه‌ای وجود دارد [۴۲].

**قضیه وجود نگاشت شبکه عصبی پس‌انتشار<sup>۵</sup> کلموگورف:** برای هر تابع  $f: [0 \ 1]^n \rightarrow \mathbb{R}^m$  که  $f(x^{(t)}, n^{(t)})$  که  $L_2$  باشد یعنی انتگرال  $\int_{[0 \ 1]^n} |f(x^{(t)}, n^{(t)})|^2 dx$  موجود باشد، یک شبکه عصبی پس-انتشار سه لایه وجود دارد که می‌تواند تابع  $f$  را با دقت میانگین مربع خطای یک عدد کسری ثابت مثبت به نام  $0 < e < d$  تقریب بزند.

فضای  $L_2$  شامل هر تابعی که در یک مسئله عملی وجود دارد می‌شود. به عنوان نمونه شامل تمام توابع ناپیوسته که به صورت تکه‌ای نیز پیوسته هستند می‌شود. این قضیه ثابت می‌کند که همیشه سه لایه کافیست. هر چند در عمل استفاده از ۴ یا ۵ و یا حتی لایه‌های بیشتر مرسوم است. زیرا بسیاری از مسائل برای حل؛ توسط شبکه سه لایه‌ای، نیازمند تعداد زیادی نرون در لایه پنهان خود می‌باشند. که این امر به کند شدن سرعت پردازش منجر می‌شود. بالاخره هر چند این قضیه تضمین می‌کند که شبکه‌های عصبی چند لایه با وزن‌های صحیح، قادر به نگاشت

<sup>2</sup> Mapping

<sup>3</sup> Feed Forward Neural Network (FFNN)

<sup>4</sup> Feed forward three layers neural network

<sup>5</sup> Back Propagation Neural Network (BPNN)

<sup>1</sup> Kolmogorov



- Transactions on Fuzzy Systems, vol. 16, no. 6, pp. 1411-1424, December 2008.
- [8] S. Guohua, L. Mingfeng, L. Teik C, "Enhanced Filtered-X Least Mean M-estimate Algorithm for Active Impulsive Noise Control," Applied Acoustics, 90: pp. 31-41, 2015.
- [9] B. Santosh Kumar, D. Debi Prasad, S. Bidyadhar, "Functional Link Artificial Neural Network Applied to Active Noise Control of a Mixture from Tonal and Chaotic Noise," Applied Soft Computing, 23: pp. 51-60, 2014.
- [10] M. Ahmadih Khanesar, E. Kayacan, M. Teshnehlal, and O. Kaynak, "Analysis of the Noise Reduction Property of Type-2 Fuzzy Logic Systems Using a Novel Type-2 Membership Function," IEEE Transactions on Fuzzy Systems, Man, and Cybernetics-part B: Cybernetics, vol. 41, no. 5, October 2011.
- [11] M. Ahmadih Khanesar, M. Teshnehlal, E. Kayacan, and O. Kaynak, "A Novel Type-2 Fuzzy Membership Function: Application to the Prediction of Noisy Data," In Computational Intelligence for Measurement Systems and Applications (CIMS), 2010 IEEE International Conference on, pp. 128-133. IEEE, 2010.
- [12] B. Anja, et al, "Wavelet Based Noise Reduction in CT-images Using Correlation Analysis," Medical Imaging, IEEE Transactions on, 27.12, pp. 1685-1703, 2008.
- [13] S. Stefan, et al, "Fuzzy Random Impulse Noise Reduction Method," Fuzzy Sets and Systems, 158.3, pp. 270-283, 2007.
- [14] S. Stefan, et al, "A Fuzzy Impulse Noise Detection and Reduction Method," Image Processing, IEEE Transactions on, 15.5, pp. 1153-1162, 2006.
- [15] T. Li; J. Jean, "Adaptive Volterra Filters for Active Control of Nonlinear Noise Processes," Signal Processing, IEEE Transactions on, 49.8, pp. 1667-1676, 2001.
- [16] L. Chien-Cheng, Y.-C. Chiang, C.-Y. Shih and C.-L. Tsai, "Noisy Time-series Prediction Using M-estimator based Robust Radial Basis Function Neural Networks with Growing and Pruning Techniques," Expert Systems with Applications 36, no. 3, 4717-4724, 2009.
- [17] Y. Y. Yao, "Human-inspired Granular Computing," Novel Developments in Granular Computing: Applications for Advanced Human Reasoning and Soft Computation, pp. 1-15, 2010.
- [18] Z. Meciaraova, "Granular Computing and its Applications in RBF with Cloud Activation Function," Journal of Information, Control and Management Systems, vol. 7, Issue 1, pp. 7-16, 2009.
- و تقریب هر تابع  $L_2$  دلخواه است؛ ولی در مورد رسیدن و یا نرسیدن به این وزن‌ها، هیچ قانون یادگیری تضمینی موجود نیست [۴۲].
- پس وقتی یک شبکه عصبی پس‌انتشار سه لایه مانند شبکه عصبی RBF گرانیولی با یک نرون در لایه میانی و آموزش پارامترها به روش «گرادیان نزولی» بتواند با تقریب بهتری نسبت به یک شبکه عصبی RBF با همان تعداد نرون در لایه میانی و همان روش آموزش؛ یک تابع غیرخطی یا آشوب را تحقق ببخشد، به طور حتم با چند نرون در لایه میانی نیز این برتری به وقوع خواهد پیوست. زیرا تابع غیرخطی و تابع آشوب به کار رفته؛ هر دو  $L_2$  می‌باشند.
- از طرفی دو قضیه «وجود نگاشت شبکه عصبی پیشرو کلموگروف» و «وجود نگاشت شبکه عصبی پس‌انتشار کلموگروف» محدود به روش آموزش خاصی نیستند؛ پس در تعمیم از دو نرون به چند نرون در لایه میانی؛ می‌توان نتیجه گرفت که یک شبکه عصبی پس‌انتشار سه لایه مانند شبکه عصبی RBF گرانیولی با آموزش پارامترها به روش «الگوریتم خوشه‌بندی K-Means» بتواند با تقریب بهتری نسبت به یک شبکه عصبی RBF با همان تعداد نرون در لایه میانی رفتار نماید.

## مراجع

- [1] X. Wu, Y. Wang, W. Liu, Z. Zhu, and Y. Tan, "Application of Nonlinear Filtering Trained RBF Networks to Multi-step Prediction of Time-series with Delayed Observations," Journal of Information & Computational Science vol. 7: 3, pp. 385-393, 2011. Available at <http://www.joics.com>.
- [2] M. Bahita, and K. Belarbi, "Neural Feedback Linearization Adaptive Control for Affine Nonlinear Systems Based on Neural Network Estimator," Serbian Journal of Electrical Engineering vol. 8, no. 3, pp. 307-323, November 2011.
- [3] D. R. Brillinger, "Time Series: Data Analysis and Theory," vol. 36, Siam, 2001.
- [4] Y. Yao, "Interval Sets and Interval-set Algebras," In Cognitive Informatics, ICCI'09, 8th IEEE International Conference on, pp. 307-314, 2009.
- [5] X. Gao, "Some properties of continuous uncertain measure," International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems, vol. 17, no. 3, pp. 419-426, 2009.
- [6] C.-F. Juang, R.-B. Huang and Y.-Y. Lin, "A recurrent self-evolving interval type-2 fuzzy neural network for dynamic system processing," Fuzzy Systems, IEEE Transactions on 17, no. 5, pp. 1092-1105, 2009.
- [7] C.-F. Juang, and Y. W. Tsao, "A self-evolving interval type-2 fuzzy neural network with online structure and parameter learning," IEEE

- [31] B. Henon, "A Two Dimensional Mapping with a Strange Attractor," Communications in Mathematical Physics, vol. 50, issue 1, pp. 69-77, 1976.
- [32] E. N. Lorenz, "Deterministic Nonperiodic Flow," Journal of the Atmospheric Sciences 20, no. 2, pp. 130-141, 1963.
- [33] M. Teshnehlal, and K. Watanabe, "Intelligent Control Based on Flexible Neural Networks," Kluwer Academic publisher, 1999.
- [34] A. Lapedes, and R. Farber, "Nonlinear Signal Processing Using Neural Networks," IEEE Conference on Neural Information Processing System, Natural and Synthetic. Proc., Mag., pp. 422, 1987.
- [35] D. E. Rumelhart, G. E. Hinton, and R. J. Williams, "Learning Internal Representations by Back Propagation Errors," Nature, 323, pp. 533-536, 1986.
- [36] V. Seydi Ghomsheh, M. Aliyari Shoorehdeli, and M. Teshnehlal, "Training ANFIS structure with modified PSO algorithm," 15th Mediterranean Conference on Control & Automation, July 27-29, 2007, Athens-Greece.
- [37] D. Kukolj and E. Levi, "Identification of complex systems based on neural and Takagi-Sugeno fuzzy model," IEEE Trans. Syst., Man, Cybern. B, Cybern., vol. 34, no. 1, pp. 272-282, Feb. 2004
- [38] V. Kreinovich, "From Interval and Probabilistic Granules to Granules of Higher Order," 2010.
- [39] P. Billingsly, "Probability and Measure," John Wiley & Sons, 2008.
- [40] S. Resnick, "Adventures in stochastic processes," Birkhäuser, Boston, 1992.
- [41] A. Papoulis, and S. U. Pillai, "Probability, Random Variables and Stochastic Processes," Tata McGraw-Hill Education, 2002.
- [۴۲] خاکی‌صدیق، علی (۱۳۸۳). ارزیابی روش‌های پیش‌بینی قیمت سهام و ارائه مدل بهینه. پژوهشکده پولی و بانکی بانک مرکزی جمهوری اسلامی ایران.
- [۴۳] علیاری شوره دلی، مهدی، محمد تشنه لب و علی خاکی صدیق (۱۳۸۳). پیش‌بینی کوتاه مدت آلودگی هوا با کمک شبکه های عصبی پرسپترون چندلایه و خط حافظه دار تاخیر. ششمین کنفرانس سراسری سیستم های هوشمند، کرمان، دانشگاه شهید باهنر کرمان. [http://www.civilica.com/Paper-ICS06-ICS06\\_025.html](http://www.civilica.com/Paper-ICS06-ICS06_025.html)
- [۴۴] ظهوری زنگنه، بختیار و عباس شولائی (۱۳۷۸). شبیه‌ساز (سیمولاتور) شبکه‌های HVDC/AC با گام بهینه. پانزدهمین کنفرانس بین المللی برق، تهران، شرکت توانیر، پژوهشگاه نیرو، [http://www.civilica.com/Paper-PSC15-PSC15\\_067.html](http://www.civilica.com/Paper-PSC15-PSC15_067.html)
- [19] Y. Yao, and Z. Ning, "Granular Computing," Wiley Encyclopedia of Computer Science and Engineering (2008).
- [20] S. Dick, A. Tappenden, C. Badke, and O. Olarewaju, "A Novel Granular Neural Network Architecture," In Fuzzy Information Processing Society, NAFIPS'07. Annual Meeting of the North American, pp. 42-47, IEEE, 2007.
- [21] B. Andrzej, and W. Pedrycz, "The Roots of Granular Computing," In GrC, pp. 806-809, 2006.
- [22] J. Abdi, B. Moshiri, and A. Khaki-Sedigh, "Comparison of RBF and MLP Neural Networks in Short-term Traffic Flow Forecasting," In Power, Control and Embedded Systems (ICPCES), 2010 International Conference on, pp. 1-4, IEEE, 2010.
- [23] D. Marcek, M. Marcek, and J. Babel, "Granular RBF NN Approach and Statistical Methods Applied to Modeling and Forecasting High Frequency Data," International Journal of Computational Intelligence Systems, vol. 2, no. 4, pp. 353-364, December, 2009.
- [24] G. B. Huang, P. Saratchandran, and N. Sundararajan, "A Generalized Growing and Pruning RBF (GGAP-RBF) Neural Network for Function Approximation," IEEE Transactions on Neural Networks, vol. 16, no. 1, pp. 57-67, January 2005.
- [25] N. Jankowski, and V. Kadirkamanathan, "Statistical Control of Growing and Pruning in RBF-Like Neural Networks," In Third Conference on Neural Networks and Their Applications, pp. 663-670, 1997.
- [26] J.-S. Roger-Jang, and C.-T. Sun, "Functional Equivalence between Radial Basis Function Networks and Fuzzy Inference Systems," IEEE Transactions on Neural Networks, vol. 4, no. 1, pp. 156-159, January 1993.
- [27] C. H. Skiadas, and C. Skiadas, "Chaotic Modeling and Simulation: Analysis of Chaotic Models, Attractors and Forms," Chapman and Hall/CRC Press is an imprint of Taylor and Francis Group, LLC, 2009, ISBN: 978-1-4200-7900-5 (hardcover: alk. paper).
- [28] J. Principe and J.-M. Kou, "Dynamic Modeling of Chaotic Time-series with Neural Networks," Advances in Neural Information Processing Systems, pp. 311-318, 1995.
- [29] V. Isham, "Statistical Aspects of Chaos, A Review," in: Networks and Chaos-Statistical and Probabilistic Aspects," pp. 124-200, Springer US, 1993.
- [30] M. C. Mackey, and L. Glass, "Oscillation and Chaos in Physiological Control System," Science 197, no. 4300, pp. 287-289, 1977.