

Xianbin SUN, Xinming JIA, Yi ZHENG, Zhen WANG, 2021. A data-driven method for estimating the target position of low-frequency sound sources in shallow seas. *Frontiers of Information Technology & Electronic Engineering*, 22(7):1020-1030.
<https://doi.org/10.1631/FITEE.2000181>

A data-driven method for estimating the target position of low-frequency sound sources in shallow seas

Key words: Vector hydrophone; Shallow sea; Low frequency; Location estimation; Recurrent neural network

Corresponding author: Xianbin SUN

E-mail: robin_sun@qut.edu.cn

 ORCID: <https://orcid.org/0000-0002-2077-5757>

Motivation

1. Compared with a deep-sea environment, a shallow sea environment has the complex characteristics of spatio-temporal variability, reflection of the signal from the shallow sea bottom, and human offshore activities that cause aliasing of target signal, thus increasing the difficulty of identification of the sound source signal.
2. With the development of noise reduction technology and stealth technology, the working frequency band of ship sonar is moving towards lower frequencies.
3. Traditional mathematical models need to set more environmental variables.

Main idea

1. Vector hydrophones can obtain information of sound pressure and acoustic particle velocity in the ocean sound field at a particular time and point.
2. As a kind of data-driven model, neural networks have the advantages of self-adaptation, self-organization, and self-learning.
3. GeoHash is a positioning method that converts two-dimensional (2D) labels into one-dimensional (1D) strings.

Method

1. A compressed recurrent neural network (C-RNN) model is proposed that compresses the signal received by a vector hydrophone into a dynamic sound intensity signal and compresses the target position of the sound source into a GeoHash code.
2. Two types of compressed data are used to carry out prior training on the recurrent neural network, and the trained network is subsequently used to estimate the target position of the sound source.

Major results

Compressed recurrent neural network

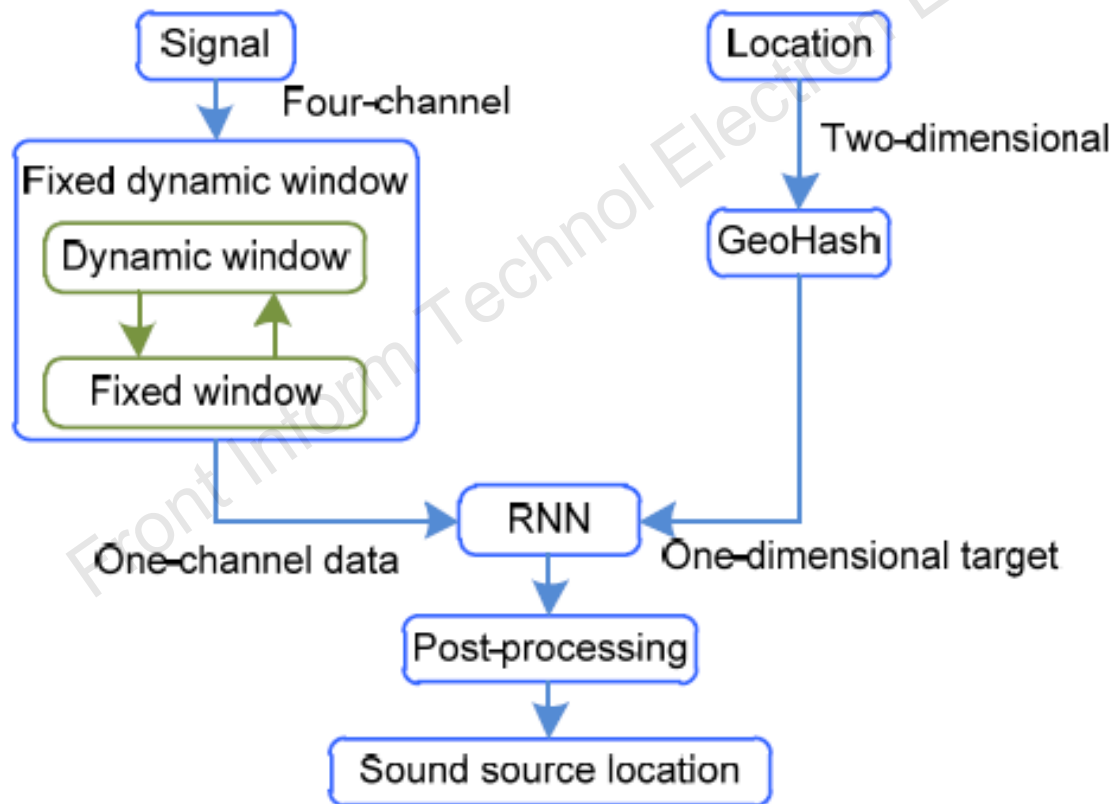


Fig. 6 Compressed recurrent neural network model

Major results (Cont'd)

Test results of our model and related methods

Table 4 Experimental results of different sound source localization models

Experiment		Lat (m)				Lng (m)				Iteration time (s)
		Before		After		Before		After		
Type	Structure	Mean	RMSE	Mean	RMSE	Mean	RMSE	Mean	RMSE	
Interception method	NW	116.0	135.7	116.0	135.7	487.4	584.9	487.4	584.9	0.04
	FW	36.3	47.2	19.7	23.8	250.7	330.3	209.1	228.8	0.53
	O-FW	25.9	35.6	14.7	18.4	154.2	240.2	124.8	146.8	0.26
	DW	130.6	155.2	130.6	155.2	304.2	365.7	304.2	365.7	0.17
	FDW	18.4	21.7	18.4	21.7	151.3	180.8	151.3	180.8	0.34
	O-FDW	17.1	23.3	8.9	10.7	104.8	177.5	57.7	70.8	0.15
Information fusion	SP	19.7	28.0	10.4	15.0	124.7	208.3	92.7	112.5	0.16
	VV	25.5	36.1	14.5	17.5	140.4	221.6	103.8	116.1	0.13
	SPVV	26.8	37.6	15.8	19.7	175.6	270.6	148.0	175.9	0.15
	SI	8.9	10.7	8.9	10.7	57.7	70.8	57.7	70.8	0.16
Network	RNN	20.4	26.6	10.7	13.0	111.3	185.7	69.0	83.3	0.24
	LSTM	60.1	65.5	59.3	64.5	244.1	276.8	239.3	271.0	0.28
	One-BiLSTM	59.7	66.2	59.2	65.0	242.2	273.2	238.7	266.6	0.50
	Two-BiLSTM	57.9	63.3	57.9	63.3	226.4	260.7	226.4	260.7	0.73
Structure	SPVV+LSTM	10.3	12.0	10.3	12.0	89.8	99.4	89.8	99.4	0.08
	SPVV+RNN	24.1	29.0	24.1	29.0	133.7	171.4	133.7	171.4	0.21
	SPVV+O+RNN	15.5	19.2	15.2	18.6	115.4	135.8	111.5	129.0	0.60
	SI+O+RNN	14.2	16.7	14.2	16.7	106.5	123.5	106.5	123.5	0.63
	GEO+SI+O+RNN	17.1	23.3	8.9	10.7	104.8	177.5	57.7	70.8	0.16
Sample	10-min+50-sample	14.6	19.5	8.7	9.9	76.1	143.5	56.4	68.1	0.20
	30-min+50-sample	63.2	90.4	48.2	57.8	223.1	330.4	145.3	162.9	0.15
	30-min+200-sample	53.6	79.4	37.5	48.8	193.5	302.8	131.5	184.6	0.18
	60-min+500-sample	35.9	44.5	35.9	44.5	188.4	300.0	188.4	300.0	0.15
	60-min+1000-sample	46.6	66.0	29.1	38.6	229.2	365.0	167.6	231.5	0.17

Before: before the averaging; After: after the averaging. Best results are in bold

Conclusions

1. A C-RNN model has been proposed which is used to achieve the localization of unknown signals.
2. The model considers that the chirp signal received by the vector hydrophone must be related to the signal in a certain time domain.
3. The C-RNN model improves the operating speed for low-frequency sound source signals, and can provide accurate positioning. The average error radius is approximately 56 m.



Xianbin SUN is an associate professor at Qingdao University of Technology. He received his BS degree from Qingdao Institute of Architecture and Engineering in 2002, his MS degree from Beijing Institute of Technology in 2006, and his PhD degree from Qingdao University of Technology. His research interests include signal processing and neural networks.



Xinming JIA received his BS degree from Northeast Petroleum University in 2018. Currently, he is an MS candidate in Qingdao University of Technology. His main research interest is neural networks.