

ISSN : 1549-523-X
UGC-CARE Listed Journal

वर्ष : 18, अंक 3-4, जुलाई-दिसम्बर 2020

Vol : 18, Issue 3-4, July-December 2020

विज्ञान प्रकाश VIGYAN PRAKASH

विज्ञान एवं प्रौद्योगिकी रिसर्च जर्नल
Research Journal of Science & Technology



लोक विज्ञान परिषद, दिल्ली
एवं
विश्व हिन्दी न्यास, न्यूयॉर्क
का प्रकाशन

विज्ञान प्रकाश - विज्ञान एवं प्रौद्योगिकी रिसर्च जर्नल
VIGYAN PRAKASH : Research Journal of Science & Technology

वर्ष : 18, अंक 3-4, जुलाई-दिसम्बर 2020

Vol : 18, Issue 3-4, July-December 2020

संस्थापक मुख्य सम्पादक / Founder Chief Editor

- स्व. प्रो. राम चौधरी / Late Prof. Ram Chaudhari
54, Perry Hill Road, Oswego, NY, 13126, USA

मुख्य सम्पादक / Chief Editor

- प्रो. ओम विकास / Prof. Om Vikas
Formerly, Dir., ABV-IIITM Gwalior; & Couns. (S&T),
Indian Embassy, Japan; & Sr. Director, Ministry of Elec. & IT
President, Lok Vigyan Parishad;
Hon. Advisor, Bharatiya Vidya Bhavan, Delhi
dr.omvikas@gmail.com

सलाहकार मण्डल / Advisory Board

- प्रो. जगदीश नारायण / Prof. Jagdish Narayan
Distinguished Chair Professor, Director, NSF Center
for Advanced Materials and Smart Structures, Dept.
of Materials Science and Engineering, Centennial Campus,
North Carolina State University, Raleigh, NC 27695-7907.
J_Narayan@ncsu.edu
- डॉ. श्याम कुमार शुक्ल / Dr. Shyam K. Shukla
Executive Director, World Hindi Foundation
44949 Cougar Circle, Fremont, CA 94539, USA
shuklas@comcast.net
- प्रो. आलोक कुमार / Prof. Alok Kumar
Department of Physics, State University of New York,
Oswego, New York 13126
Alok.kumar@oswego.edu
- प्रो. आदित्य शास्त्री / Prof. Aditya Shastri
Vice Chancellor, Banasthali University, Rajasthan
saditya@banasthali.ac.in

प्रकाशन सहयोग / Publication Support

लोक विज्ञान परिषद् / Lok Vigyan Parishad
C-15, Tarang Apartment, 19 IP Extension, Delhi-110092
www.LokVigyanParishad.com

सम्पादक मण्डल / Editorial Board

- प्रो. ओउम प्रकाश शर्मा / Prof. Oum Prakash Sharma
Director, NCIDE, IGNOU, New Delhi-110068
Gen. Secy., Lok Vigyan Parishad
opsharma@ignou.ac.in
- प्रो. अनुपम शुक्ल / Prof. Anupam Shukla
Director, Indian Institute of IT Pune, Maharashtra
dranupamshukla@gmail.com
- प्रो. (डॉ.) अजय चौधरी / Prof. (Dr.) Ajay Choudhary
Professor & Head, Dept. of Neurosurgery,
ABV-Institute of Medical Sciences
RML Hospital, New Delhi
ajay7.choudhary@gmail.com
- प्रो. प्रतापानन्द झा / Prof. Pratapanand Jha
Director, Cultural Informatics Lab (CIL), IGNC, New Delhi
pjha@ignca.nic.in
- प्रो. कृष्ण कुमार मिश्र / Prof. Krishna Kumar Mishra
Homi Bhabha Centre for Science Education, TIFR,
Mumbai - 400088, (India)
kkm@hbcse.tifr.res.in
- श्री रामशरण दास / Mr. Ram Sharan Das
49, Sector-4, Vaishali, Ghaziabad 201010, U.P.
rsgupta_248@yahoo.co.in

ऑनलाइन प्रदर्श (वैबसाइट) / Online Presence (Website)

- दिव्या शर्मा / Divya Sharma
Designer's Bliss, Sydney, NSW, Australia
www.designersbliss.com

मुद्रक / Printer

रोहित कौशिक / Rohit Kaushik
arthstudios@yahoo.co.in

विषय क्रम

- संपादकीय : लोकल-ग्लोबल साधती भारत शिक्षा नीति : कतिपय सतर्कता बिन्दु /
Editorial : NEP2020 Balancing Local and Global : a few points of caution 2
ओम विकास

शोध आलेख / Research Articles

- मशीन लर्निंग द्वारा डायबिटीज डिटेक्शन- वर्गीकरण तकनीकों का एक तुलनात्मक अध्ययन 6
/ Machine Learning Based Diabetes Detection-A Comparative Study of Classifiers
सोफिया गोएल एवं डा० सुधांशु शर्मा
- नैचुरल लैंग्वेज प्रोसेसिंग में ऊर्दू स्टेमर का विकास /Development of Urdu Stemmer in 18
Natural Language Processing
वैशाली गुप्ता, निशीथ जोशी एवं इति माथुर
- वायुमंडलीय कणिकीय तत्वों PM2.5 एवं PM10 के स्तर पर लॉकडाउन का प्रभाव 29
नई दिल्ली के सन्दर्भ में / Effect of lockdown on the levels of atmospheric
particulate matter PM2.5 and PM10: A case study from New Delhi
शिल्पी शाक्य एवं बिन्ध्या चल यादव

भारत ज्ञान परंपरा / Indian Knowledge Tradition

- भारतीय उद्योग के पितामह - जमशेद जी टाटा / Father of Indian Industry – 39
Jamshedji Tata
वीरेंद्र ग़ोवर
- पुस्तक समीक्षा / Book Review 44

प्रतिक्रियाएं / Feedback 48

विज्ञान प्रकाश रिसर्च जर्नल में प्रकाशित लेख/सामग्री लेखकों के अपने निजी विचार हैं।

विज्ञान प्रकाश के संपादक मंडल तथा प्रकाशक का कोई दायित्व नहीं है।

सम्पादकीय

लोकल-ग्लोबल साधती भारत शिक्षा नीति : कतिपय सतर्कता बिन्दु

NEP2020 Balancing Local and Global : Few Points of Caution

भारत स्वतंत्र हुआ। औपनिवेशिकता के खुमार को कम करने के लिए अपनी देशज शिक्षा नीति बनाने के प्रारम्भिक प्रयास में प्रमुख हैं :

पहला, राधाकृष्णन उच्चशिक्षा कमीशन (1948) जिसमें ज्ञान-विज्ञान, राजनीति, उद्योग, आदि क्षेत्रों में नेतृत्व और कौशल विकास पर बल दिया गया। सिफारिश में उच्च शिक्षा में नैतिक मूल्यों का संवर्धन, कुछ मिनट का ध्यान, भारतीय भाषा माध्यम, धर्म का दार्शनिक अध्ययन, आदि शामिल थे।

दूसरा, शिक्षा पर राष्ट्र नीति (National Policy on Education) : NPE 1968 में और इसका परिष्कार 1992 में, तकनीकी एवं व्यावसायिक शिक्षा के लिए प्रवेश परीक्षा में एकरूपता, हिन्दी-अंग्रेजी और क्षेत्रीय भाषा को अपनाना, संस्कृत भाषा के अध्ययन को बढ़ावा देना और शिक्षा पर GDP का 6% व्यय करना इसकी महत्वपूर्ण सिफारिशें थीं।

तीसरा, राष्ट्रीय ज्ञान आयोग (National Knowledge Commission): 2007 में 21वीं सदी की चुनौतियों के परिप्रेक्ष्य में शिक्षा तंत्र को उत्कृष्ट बनाने की दिशा में विज्ञान-प्रौद्योगिकी संस्थानों में ज्ञान-सृजन को बढ़ावा देने के लिए सुझाव दिए गए। 27 विषय क्षेत्रों को इंगित किया गया जैसे भाषा, अनुवाद, ग्रंथालय, ज्ञान-नेट, स्कूली शिक्षा, बौद्धिक सम्पदा अधिकार, रिसर्च, उद्यमिता आदि। इसमें तात्कालिक वैश्विक प्रतिद्वंद्विता पर जोर था। स्कूली शिक्षा में अंग्रेजी पर बल दिया गया। कक्षा 1 से अंग्रेजी की शिक्षा दिए जाने की वकालत की गई।

राधाकृष्णन कमीशन (1948) की सिफारिशें भारतीय संस्कृति प्रधान थी, लेकिन इसमें केवल उच्च शिक्षा की बात की गई थी, NPE (1968/92) की सिफारिशें कार्यान्वयन में सुगमतापरक थीं जबकि 2007 की NKC (राष्ट्रीय ज्ञान आयोग) सिफारिशें औद्योगिक दृष्टि से मानव संसाधन तैयार करने की थीं।

चौथा, राष्ट्रीय शिक्षा नीति (National Education Policy) 2020 नई सोच से प्रेरित है। भारतीय संस्कृति से सम्पन्न उद्यमशील प्रतिभा विकास और उत्कृष्ट नवाचार परक है। NEP 2020 में चार अध्याय हैं – स्कूली शिक्षा (48 पृष्ठ), उच्चतर शिक्षा (48 पृष्ठ), अन्य विचारणीय मुद्दे (15 पृष्ठ), और क्रियान्वयन रणनीति (15 पृष्ठ)।

पहले अध्याय में वर्णित स्कूल शिक्षा नीति में इस बात का ध्यान रखा गया है कि 6 वर्ष तक के बच्चों के मस्तिष्क का 85% विकास हो जाता है। इसलिए 3 से 6 तक आयुवर्ग के लिए बाल शिक्षा में पोषण, स्वास्थ्य और आचरण पर बल दिया गया तथा आयु 7 से 8 को कक्षा 1 व 2 के समान रखा गया। आयु 3 से 8 काल को व्यक्ति और राष्ट्र की ज्ञान-क्षमता एवं संस्कार की नींव की भांति देखा जाए, यह माना गया। प्राथमिक शिक्षा 5 वर्षों की जगह 8 वर्षों की जगह प्रस्तावित है – बुनियादी शिक्षा 3-8 आयु में और तैयारी शिक्षा काल 9-11 आयु में विभाजित किया गया। शेष वर्गीकरण पूर्ववत् रहा। इस प्रकार स्कूली शिक्षा 5+3+3+4 चरणों में रखी गई है।

बुनियादी शिक्षा की सफलता राज्य सरकारों पर निर्भर है। समग्र अधिगम और एक पुस्तक की नीति के आधार पर सर्व शिक्षा अभियान को सफल बनाया जाए। शक्यता (Can do) और संकल्प (Will do) का भाव दृढ़ हो। शारीरिक और मानसिक स्वास्थ्य के साथ नवोन्मेषी अनुसंधानपरक प्रवृत्ति का विकास हो। कक्षा 6-12

तक कौशल विकास पर बल दिया जाए। संस्कृत में विज्ञान और भाषा की वैज्ञानिकता को भी बताया जाए। आर्टिफिशियल इंटेलिजेंस (AI) और कोडिंग को शामिल करना अच्छा है, लेकिन इसे भारतीय भाषा में भी समझाया और अभ्यास कराया जाए, कुंजियों पर नागरी वर्ण उकैरे हुए फ्लैक्स्ज कीबोर्ड अनिवार्यतः दिए जाएं।

उद्योग कार्यबल (19–24 वर्ष के युवा) वोकेशन शिक्षा के माध्यम से भारत में 5%, और दक्षिण कोरिया में 96% है। विडम्बना है कि पिछले चार दशक के प्रयासों के बावजूद भारतीय समाज में वोकेशनल शिक्षा को दायम दर्जे का मानते हैं, प्रतिष्ठित नहीं बन सकी। सुझाव है कि एक, एकेडेमिक शिक्षा में कौशल विकास पर करीब 40% बल दिया जाए। दो, वोकेशन सर्टिफिकेट को एकेडेमिक सर्टिफिकेट के समतुल्य माना जाए जिससे उच्चशिक्षा में प्रवेश आसान हो। तीन, क्रेडिट ट्रांसफर की व्यवस्था हो। किसी अन्य मान्यता प्राप्त संस्थान से ट्रेनिंग को मान्यता दी जाए।

दूसरे अध्याय में उच्च शिक्षा के सम्बंध में सिफारिशें हैं। UG व PG स्तर की शिक्षा में भारतीय संस्कृति परिचय, विविध विषयों का समुचित समावेश, विशेषज्ञता, कौशल-क्षमता, सहकारिता, नैतिकता और नवोन्वेषी अनुसंधान परक प्रवृत्ति के विकास पर बल दिया जाए। सभी प्रतिभावान प्रतिभागियों को गुणवत्तापूर्ण उच्च शिक्षा सुलभ हो। प्रवेश और प्रोग्राम छोड़ने का लचीला प्रावधान हो। अन्य संस्थान से कोर्स करने पर क्रेडिट मान्य हों।

उच्च शिक्षा में GER प्रवेश अनुपात 26.3% (2018) से 50% (2035) बढ़ाने का लक्ष्य है। कोरोना काल में ऑन-लाइन क्लास की सहज स्वीकृति हो गई। लेकिन ऑन-लाइन और इन क्लास मिश्रित शिक्षण विधा (Blended hearing) प्रभावी होगी। UG स्तर पर बीच में भी छोड़ने और बाद में फिर आगे पढ़ाई करने का भी प्रावधान किया है। 1 वर्ष के बाद छोड़ने पर सर्टिफिकेट, 2 वर्ष बाद डिप्लोमा, 3 वर्ष बाद एडवांस्ड डिप्लोमा ही, डिग्री 4 वर्ष बाद ही दी जाए। 3 वर्ष के बाद डिग्री दिए जाने की सिफारिश उचित नहीं है।

Washington Accord के अनुसार 4 वर्षीय डिग्री प्रोग्राम मान्य है। NBA से accreditation भी 4 वर्ष के डिग्री प्रोग्राम को मिल सकेगा। हरेक विश्वविद्यालय/संस्थान को समग्र पाठ्यक्रम (holistic curriculum) बनाने, जरूरी संसाधन (Resources) जुटाने और लक्षित परिणाम (Outcome) के बारे में स्पष्ट योजना बनाने की आवश्यकता है।

समग्र पाठ्यक्रम (holistic curriculum) में अन्य विषयों को उपयोगिता के आधार पर जोड़ा जाए। NEP 2020 में डिग्री प्रोग्राम से अपेक्षाएं अधिक हैं, किसी एक विषय में गहराई और अन्य रुचिकर विषयों का सतही ज्ञान हो। हिन्दी/भारतीय भाषा को युक्तिपूर्ण तरीके से उच्चतर शिक्षा का माध्यम बनाया जाए। *संक्रमण काल में मिश्रित विधा को अपनाएं जिसमें बोलचाल की भाषा में बात समझायें, लेकिन प्रचलित तकनीकी शब्द अंग्रेजी के भी रहें। यदि किसी तकनीकी शब्द का हिन्दी समतुल्य शब्द है तो उसका प्रयोग भी साथ-साथ करें।* हिन्दी शब्द सामाजिक संवेदना का भाव देता है, तो अंग्रेजी का प्रचलित तकनीकी शब्द वैश्विक (global) प्रतिद्वंद्विता के प्रति तैयारी का प्रतीक है। परिवेश (local) से जुड़े रहें और वैश्विक (global) प्रतिद्वंद्विता भी बनें। इससे इंडस्ट्री में जॉब पाने में भी परेशानी नहीं होगी।

हिन्दी/भारतीय भाषा माध्यम में तकनीकी शिक्षा के लिए क्या हम तैयार हैं ?

अभी नहीं, आज तकनीकी शब्दावली वेब पर खोजी नहीं जा सकती, सर्चबल नहीं है। उच्चशिक्षा में भारतीय भाषा में शब्द निर्माण एवं निर्वचन और तकनीकी अनुवाद विषयों को पढ़ाया नहीं जाता। इसलिए

वैज्ञानिक नए शब्द सृजन में असहाय रहते हैं। तकनीकी पुस्तकें हिन्दी/भारतीय भाषा में लाने के लिए प्रस्तावित एक अनुवाद एवं निर्वचन संस्थान (ITI) बनाने से ही काम नहीं चलेगा।

सुझाव है कि तकनीकी संस्थानों के समूह बनें और विषयवार अनुसृजन कार्य किया जाए। पढ़ाने के बाद शब्दावली प्रणयन गुणवत्ता पूर्ण होगा। तकनीकी शिक्षा में शिक्षण-विधाओं (Pedagogy) एवं अनुसृजन से परिचित कराने के लिए शिक्षकों के प्रशिक्षण का प्रावधान हो। यह 2-3 माह का कोर्स हो सकता है। ग्रीष्म काल के अवकाश में ऐसा कोर्स चला सकते हैं। NITTR में इस प्रकार के कोर्स की व्यवस्था हो।

तकनीकी और मेडीकल शिक्षा में हिन्दी/भारतीय भाषा का प्रयोग सामाजिक अपेक्षा है, जिससे प्रशिक्षित विशेषज्ञ अपने परिवेश और देश की समस्या को उत्कृष्टता के साथ हल कर सकें। लौकिक परिवेश में पारस्परिक विचार विनिमय में आसानी हो, और इसके साथ वैश्विक स्तर पर भी प्रतिद्वंदी बनें। इसलिए आवश्यक है कि उन्हें हिन्दी/भारतीय भाषा और अंग्रेजी का भी अच्छा ज्ञान हो, दोनों में काम करने, और समझने समझाने की क्षमता हो।

चुनौतीपूर्ण पाठ्यक्रम निर्माण एवं कार्यान्वयन :

पाठ्यक्रम (Curriculum) NEP2020 के अनुरूप समग्र (Holistic) हो, सुचारु चलाने के लिए समुचित संसाधन (Resources) हों। शिक्षण संस्थान का लक्ष्य रहेगा कि संस्थान की प्रतिष्ठा में बढ़ोतरी हो। संस्थान की प्रतिष्ठा का मापदण्ड है छात्रों को अच्छी नियुक्तियाँ मिलें, अधिकतम प्लेसमेंट हो, शिक्षकों को रिसर्च और प्रमोशन के अवसर मिलें, और शिक्षण संस्थान को उच्च, गुणवत्तापूर्ण संस्थान के रूप में पहचान मिले। तभी अधिक से अधिक मेधावी छात्र एडमिशन लेंगे।

शिक्षकों को सम्मान प्रमोशन की बात इस लिए कहनी पड़ रही है कि एकेडेमिक पर्फॉर्मेंस इंडेक्स (API : Academic Performance Index) की गणना करते समय शिक्षण के द्वारा हिन्दी/भारतीय भाषा में पढ़ाने, नोट्स बनाने, प्रोजेक्ट रिपोर्ट लिखाने और रिसर्च पेपर लिखने को हेय न समझा जाए, प्रत्युत इसे वांछनीय मानकर 20 प्रतिशत रिसर्च पेपर हिन्दी/भारतीय भाषा में आवश्यक किए जाएं। *NAAC, NBA और NIRF में मूल्यांकन पैरामीटर में बदल करना अनिवार्य है, अन्यथा शिक्षा नीति का सफल कार्यान्वयन सम्भव नहीं है।*

पाठ्यक्रम तैयार करना चुनौतीपूर्ण हो सकता है, क्योंकि डिग्री की समयावधि 4 वर्ष में संस्कृत भाषा परिचय, भारतीय संस्कृति, भारत ज्ञान परम्परा, तकनीकी अनुसृजन के सिद्धान्त जैसे नए विषय जुड़ेंगे। CBCS (चॉइस बेस्ड क्रेडिट सिस्टम) के अन्तर्गत कोर विषय से परे मेडीकल, म्यूजिक, भाषा, प्रबंधन आदि किसी विषय का अध्ययन भी चुन सकते हैं। थ्योरी के साथ स्वयं कर सकने के कौशल का भी विकास हो। इन्नोवेशन/रिसर्च आधारित प्रोजेक्ट भी करे। विद्यार्थी से अपेक्षाएं बहुत, लेकिन समय कम है। पाठ्यक्रम बनने के बाद सुबोध पाठ्य-सामग्री बनाना भी आवश्यक है। कई विषय, कई पुस्तकें, कई NPTEL / SWAYAM पर लेक्चर हैं, कई प्रोफेसर नई पुस्तकें लिखना चाहते हैं। भाषान्तर के लिए सर्चबल तकनीकी शब्दावली हो, पारिभाषिक शब्दावली हो, शब्द निर्माण की सहायक सामग्री संस्कृत-हिन्दी, संस्कृत-इंग्लिश, हिन्दी समानान्तर, अंग्रेजी-हिन्दी शब्दकोश मिलें। हिन्दी/भारतीय भाषा में सुबोध पाठ्यक्रम सामग्री तेजी से तैयार कराने के लिए शिक्षक लेखकों को समेकित अनुसृजन सिस्टम (ITS : Integrated Transcreation System) उपलब्ध कराए जाएं। MEITY (इलेक्ट्रॉनिक्स एवं सूचना प्रौद्योगिकी मंत्रालय) के माध्यम से ITS को उपलब्ध कराया जा सकता है। AICTE से भी प्रोत्साहन संभव है।

समेकित अनुसृजन सिस्टम (Integrated Transcreation System) में मशीन ट्रांसलेशन सिस्टम, OCR, हस्तलेख पहचान तंत्र, स्पीच टू टैक्स्ट (S2T), टेक्स्ट टू स्पीच (T2S), ट्रांसलेशन मोमोरी, सर्चबल तकनीकी शब्दावली, संस्कृत, हिन्दी, अंग्रेजी के शब्दकोश, समान्तर कोर्पोरा, संक्षेपक (Summarizer), उत्तम लेखों को संग्रह करने की सुविधा, Chrome, Firefox, MS Edge आदि ब्राउजर पर 10 सुंदर यूनिकोड फॉन्ट, नागरी वर्ण उक्रे हुए INSCRIPT इंस्क्रिप्ट मानक कीबोर्ड, आदि का समावेश हो। समेकित अनुसृजन सिस्टम से सुबोध पाठ्यसामग्री तैयार करने की गति चार पांच गुनी बढ़ जाएगी। मौलिक लेखन की प्रवृत्ति में भी संवर्धन होगा।

शोध संवर्धन : भारत का ग्लोबल इन्नोवेशन इंडेक्स 2015 में 81 था, 2019 में 52 हो गया, और 2020 में 48. इन्नोवेशन में प्रति वर्ष सुधार हो रहा है। रिसर्च के सम्बंध में दो बातें हैं— Quality (गुणवत्ता) और Relevance (सार्थकता/उपादेयता)। भारत में रिसर्च गुणवत्तापूर्ण जर्नल में प्रकाशित किए जाने के प्रयास किए जाते हैं, जिससे उसके Citation यानि दूसरों के द्वारा उसे खोलने-पढ़ने की संख्या बढ़े और एकेडेमिक परफोर्मेंस इंडेक्स (API) की गणना में लाभ मिले, प्रमोशन मिले। कई विभाग और मन्त्रालय भी शोध विकास के लिए वित्तीय सहायता देते हैं। इनके प्रोजेक्ट का लाभ प्रायः सामान्य जन तक नहीं पहुंचता। करोड़ों रुपए खर्च होते हैं, लेकिन ढाक के तीन पात। *सुझाव है कि सभी शोध विकास परियोजनाओं का रिसर्च ऑडिट अनिवार्य हो।*

प्रस्तावित NRF (नेशनल रिसर्च फाउण्डेशन) से अधिक से अधिक युवा वैज्ञानिकों को रिसर्च के लिए ग्रांट की व्यवस्था हो। आत्मनिर्भर भारत के सम्बंध में रिसर्च देश की समस्याओं को हल करने में सहायक हो, उपयोगी हो। NRF बेसिक रिसर्च पर फोकस करे, जो भविष्योन्मुखी हो, समावेशी हो, वैश्विक प्रतिद्वंद्वी हो और वैश्विक नेतृत्व प्रदान करने में सक्षम बने। इसके लिए खुलापन (openness), पारदर्शिता (Transperancy) और सहभागिता (collaboration) आवश्यक हैं। पब्लिक प्राइवेट पार्टनरशिप में रिसर्च लैब खोली जाएं, जहाँ जिज्ञासामूलक शोध (curiosity driven research) और औद्योगिक शोध-विकास पर मिलकर काम हो।

NEP 2020 में बहु विषयक विश्वविद्यालय (Multi-disciplinary University) बनाने और समग्र शिक्षा (Holistic Education) की सिफारिश की गई है। लेकिन यह लचीले ढंग से अन्य विशेषज्ञता संस्थानों के सहकार सहयोग से सम्भव है। 80-20 के अनुपात में अध्यापन प्रधान (Teaching Predominant) और रिसर्च केन्द्रित (Research Focused) स्वायत्त शिक्षण संस्थाएं चिह्नित की जाएं। इसी स्वीकृति के आधार पर इन संस्थानों को उत्कृष्ट बनने के लिए प्रोत्साहित किया जाए। प्रत्येक रिसर्च केन्द्रित संस्थान 5-7 अध्यापन प्रधान कॉलेजों को मेंटर करने की जिम्मेदारी ले। शिक्षण कौशल उत्कर्षण, छात्रों को रिसर्च लैब की सुविधा, पाठ्यक्रम, अध्यापन, नवाचारमय प्रयोगात्मक कौशल विकास आदि में सुधारात्मक सहयोग दें। उच्च शिक्षा संस्थानों के पाठ्यक्रम में संस्कृत का वैज्ञानिक परिचय और संस्कृत वाङ्मय में विज्ञान के विषय शामिल किए जाएं। तकनीकी संस्थानों में तकनीकी अनुसृजन कोर्स भी दिया जाए जिससे वे आधुनिक विज्ञान एवं प्रौद्योगिकी साहित्य का हिन्दी/भारतीय भाषा में अनुवाद, अनुसृजन एवं मौलिक लेखन में समर्थ हों। विशेषज्ञ समिति के द्वारा डिजिटल शिक्षण सामग्री की गुणवत्ता को सुनिश्चित किया जाए।

. . . ओम विकास

dr.omvikas@gmail.com

मशीन लर्निंग आधारित डायबिटीज संसूचन :
वर्गीकरणकारी तकनीकों का एक तुलनात्मक अध्ययन

**Machine Learning Based Diabetes Detection
- A Comparative Study of Classifiers**

श्रीमती सोफिया गोयल¹, डॉ. सुधांशु शर्मा²

¹रिसर्च स्कॉलर, सोसिस, इग्नू, नई दिल्ली, इंडिया

²असिस्टेंट प्रोफेसर, सोसिस, इग्नू, नई दिल्ली, इंडिया

¹sofiagoel@gmail.com; ²sudhansh@ignou.ac.in

सारांश

मधुमेह पूरी दुनिया के स्वास्थ्य के लिए एक चिंता जनक विषय है, इसका निदान और इलाज चिकित्सा क्षेत्र के लिए प्रमुख चुनौती है, क्योंकि इसे नियंत्रित किया जा सकता है लेकिन इसे ठीक नहीं किया जा सकता है। इस रोग का निदान जितना जल्दी होगा, रोगी के लिए उतना ही अच्छा होगा, और इस बीमारी के निदान के लिए हम मशीन लर्निंग (एमएल-ML) की तकनीक का उपयोग कर सकते हैं। हम जानते हैं कि मधुमेह का समय पर निदान हो जाए तो रोगी को विभिन्न प्राण घातक रोगों से बचाया जा सकता है। इस कारण, प्रस्तुत पेपर (आर्टिकल) में हम मशीन लर्निंग की विभिन्न वर्गीकरण तकनीकों का तुलनात्मक अध्ययन प्रस्तुत करने का प्रयास कर रहे हैं। मशीन लर्निंग (एमएल) में विभिन्न वर्गीकरण तकनीक उपलब्ध हैं, जैसे कि सपोर्ट वेक्टर मशीन (Support Vector Machine-SVM) (एस.वी.एम), के-नियरेस्ट नेबर (K-Nearest Neighbor-KNN) (के.एन.एन), एल्गोरिथ्म, डिसिशन ट्री (Decision Tree), आर्टिफिशियल न्यूरल नेटवर्क (Artificial Neural Network-ANN) (ए.एन.एन), नैव बैस क्लासिफायर (Naïve Bayes Classifier) आदि। लेकिन सवाल यह है कि मधुमेह की पहचान करने में वर्गीकरण तकनीकों में से कौन सी तकनीक सटीक है, क्योंकि निदान की सटीकता अधिक महत्वपूर्ण है। प्रस्तुत अध्ययन में वर्गीकरण तकनीकों की सटीकता जांचने के लिए हमने UCI रिपोजिटरी द्वारा जारी किए गए पीमा इंडियन (PIMA INDIAN) नामक डाटासेट पर डाटा इंस्पूटेशन तकनीक का उपयोग किया है। उसके बाद तुलनात्मक अध्ययन के लिए हमने इस डाटा सेट पर मशीन लर्निंग की विभिन्न वर्गीकरण तकनीकों का प्रयोग किया और मधुमेह निदान की सटीकता जाँची है। तुलनात्मक अध्ययन के फलस्वरूप यह पाया गया कि आर्टिफिशियल न्यूरल नेटवर्क तकनीक की सटीकता सर्वोत्तम, 88.6% है, जो अन्य तकनीकों के मुकाबले काफी ज्यादा है। सटीकता का यह स्तर इंगित करता है कि आर्टिफिशियल न्यूरल नेटवर्क में डेटा को वर्गीकृत करने की उच्च क्षमता है, जो इसकी हाई कंप्यूटिंग आर्किटेक्चर से संबंध रखता है, जिसके कारण बिना फीचर इंजीनियरिंग के इसको प्रोग्रामित किया जा सकता है।

Abstract

Diabetes is a matter of concern for the health of the entire world, its diagnosis and cure are among the prime challenges for the medical fraternity, because it can be controlled but can't be cured, sooner the diagnosis the better it will be for the patient. Thus, use of machine learning for

earlier classification of diabetes plays a vital role to protect patient from the life threatening complications in future, various classification techniques are available in Machine Learning (ML) viz. Support Vector Machines (SVM), K-Nearest Neighbor (KNN) algorithm, Decision Trees, Artificial Neural Networks (ANN), Naïve Bayes Classifier etc. But the question is which of the classification techniques is quite accurate in identifying the Diabetes, accuracy of diagnosis is more important. Thus, in the present work the assurance of accuracy level is strengthened by applying the data imputation techniques on the dataset i.e. dataset of Pima Indian women of Arizona named as PIMA INDIAN DATASET from UCI repository. Thereafter the said ML techniques were compared, and it is found that among the said classifiers ANN shows high performance as compared to other conventional classifiers. Highest accuracy achieved with ANN is 88.6%. This level of accuracy indicates that neural network has high potential to classify the data, which relates to its high computing layer architecture that can be programmed without any featured engineering.

मुख्य शब्द : सपोर्ट वेक्टर मशीन (एस.वी.एम), के-नियरेस्ट नेबर (के.एन.एन) एल्गोरिथ्म, डिसीशन ट्रीज, आर्टिफिशियल न्यूरल नेटवर्क्स (ए.एन.एन), नैव बैस क्लासिफायर, डेटा इंप्यूटेशन, पीमा इंडियन नामक डाटा सेट।

Keywords : Support Vector Machines (SVM), K-Nearest Neighbor (KNN) algorithm, Decision Trees, Artificial Neural Networks (ANN), Naïve Bayes Classifier, Data imputation, PIMA Indian dataset.

परिचय :

मधुमेह एक घातक बीमारी है, जो दुनिया भर में बहुत तेजी से फैल रही है। इस बीमारी में शरीर

पर्याप्त इंसुलिन का उत्पादन करने में असमर्थ हो जाता है, जिससे रक्त शर्करा (ब्लड शुगर) का स्तर बढ़ जाता है, जो आगे चलकर खतरनाक बीमारियों [2] का कारण बनता है। यह अनियंत्रित रक्त शर्करा कई जटिलताओं और विकारों का मूल कारण है, जैसे – दृष्टि विकार, हृदय संबंधी समस्याएं, धमनी उच्च रक्तचाप, त्वचा की समस्याएं, स्ट्रोक, तंत्रिका क्षति, गुर्दे की विफलता यहां तक कि दिल का दौरा भी, अनियंत्रित रक्त शर्करा से हो सकता है।

अनियंत्रित रक्त शर्करा से उत्पन्न होने वाली इन जटिलताओं को देखते हुए शोधकर्ताओं ने विभिन्न पारंपरिक मशीन लर्निंग तकनीकों का उपयोग करके, रोग को वर्गीकृत करने के लिए कई शोध किए हैं, लेकिन रोग का पता लगाने की सटीकता (एक्यूरेसी) के स्तर के संदर्भ में एक तुलनात्मक अध्ययन हमेशा समय की आवश्यकता के रूप में महसूस किया जाता रहा है। इस प्रस्तुत पत्र में हमने पारंपरिक मशीन लर्निंग की तकनीकों के साथ आर्टिफिशियल न्यूरल नेटवर्क्स (ANN) तकनीक की डायबिटीज डिटेक्शन सटीकता से तुलना की है, क्योंकि डायबिटीज डिटेक्शन की सटीकता अधिक महत्वपूर्ण है। प्रस्तुत अध्ययन में वर्गीकरण तकनीकों की सटीकता जांचने के लिए हमने UCI रिपॉजिटरी द्वारा जारी किये गए पीमा इंडियन (PIMA INDIAN) नामक डाटा सेट [1] पर डाटा इंप्यूटेशन तकनीक का उपयोग किया है। उसके बाद तुलनात्मक अध्ययन के लिए इस डाटा सेट पर हमने मशीन लर्निंग की विभिन्न वर्गीकरण तकनीकों का उपयोग किया और उन वर्गीकरण तकनीकों की मधुमेह डिटेक्शन में सटीकता की तुलनात्मक जाँच की है।

UCI रिपॉजिटरी द्वारा जारी किये गए पीमा इंडियन (PIMA INDIAN) नामक डाटा सेट [1] में 8 एट्रिब्यूट, 768 इंस्टैंस और 1 बाइनरी क्लास एट्रिब्यूट शामिल हैं। विशेषताओं का विवरण तालिका -1 में नीचे प्रस्तुत किया गया है।

तालिका 1. पीमा भारतीय डेटासेट के एट्रिब्यूट

क्र.सं.	विशेषता का नाम	विवरण
1	गर्भावस्था	संख्यात्मक (Numeric)
2	प्लाज्मा ग्लूकोज कंसंट्रेशन	ग्लूकोज टॉलरेंस टेस्ट में 2 घंटे
3	डायस्टोलिक रक्तचाप	mm Hg
4	ट्राइसेप्स स्किन फोल्ड थिकनेस	Mm
5	2- एच सीरम इंसुलिन	mu U/ml
6	बॉडी मास इंडेक्स (BMI)	Kgm-2
7	मधुमेह वंशावली समारोह	संख्यात्मक (Numeric)
8	आयु	वर्ष

इस डेटासेट के चुनाव के पीछे का मुख्य कारण यह है कि इसमें गर्भावस्था और रक्त प्लाज्मा विवरणों के साथ उम्र, अधिक वजन, बीएमआई, 2 एच सीरम इंसुलिन, डायस्टोलिक रक्तचाप जैसी विशेषताएं भी शामिल हैं। दूसरी बात यह है कि इसमें डायबिटीज पेडिग्री फंक्शन (DPF) शामिल है, जो कि रिश्तेदारों में डायबिटीज मेलिटस के इतिहास और मरीज के रिश्तेदारों के आनुवांशिक संबंधों के आंकड़ों से संबंधित है।

इसके अलावा, यह देखा गया कि कई डेटासेटों की तरह, इस डेटासेट में भी कहीं कहीं पर कुछ आंकड़ों की अनुपस्थिति थी (Missing data values), जो कि मजबूत तकनीकी वर्गीकरण मॉडल तैयार करने के लिए शोधकर्ताओं द्वारा सबसे बड़ी चुनौती है। इस लिए, हमारा शोध सबसे पहले अनुपस्थित आंकड़ों को और ऑउटलायर्स के निरीक्षण के लिए है [3], निरीक्षण के उपरांत उन आंकड़ों को सही करने के पश्चात, हमने विभिन्न वर्गीकरण तकनीकों का डेटासेट पर प्रयोग किया और उन वर्गीकरण तकनीकों की मधुमेह डिटेक्शन में सटीकता की तुलनात्मक जाँच की है।

पेपर का शेष भाग निम्नानुसार आयोजित किया गया है: खण्ड (Section)2 में साहित्य समीक्षा, खण्ड (Section)3 में मेथड्स और डेटासेट शामिल हैं,

खण्ड (Section) 4 में अनुसंधान के लिए अपनाई गई पद्धति का वर्णन किया है। इसके अलावा खण्ड (Section)5 में परिणाम और चर्चा प्रस्तुत की गयी है, और अंत में खण्ड (Section)6 में प्रस्तुत शोधपत्र का निष्कर्ष दिया गया है।

साहित्य समीक्षा :

शोधकर्ताओं ने विभिन्न पारंपरिक मशीन लर्निंग दृष्टिकोणों का उपयोग करके, रोग को वर्गीकृत करने के लिए कई शोध किए हैं, और इन शोधों के अध्ययन से यह पता चलता है कि टाइप 2 डायबिटीज मेलिटस (T2DM) का सटीक विश्लेषण करने के लिए एक मॉडल की जबरदस्त जरूरत है। नाहला एट. अल. (Nahla et al.) [5] ने, मधुमेह रोग के निदान के लिए सिक्वेन्शियल कवरींग अप्रोच (Sequential Covering Approach) लागू करके एक सपोर्ट वेक्टर मशीन (Support Vector Machine & SVM) आधारित मॉडल प्रस्तुत किया। उनके प्रस्तुत काम में डेटा की सब-सैम्पलिंग के लिए के-मीन्स क्लस्टरिंग तकनीक का उपयोग किया गया व उनके नियम रक्त शर्करा के स्तर (एफबीएस-FBS-Fasting Blood Sugar) और कमर परिधि (Waist Circumference) जैसे कारकों पर आधारित थे। इस प्रस्तावित मॉडल में, डायबिटीज की बीमारी को डायग्नोज करने के लिए, डेटा के बजाय मॉडल

से प्राप्त नियमों का उपयोग किया गया है। झेंग एट अल (Zhang et. al.) [5] ने 23,281 मधुमेह से संबंधित रोगियों में से 300 रोगियों के नमूनों पर एक अध्ययन किया, उनका प्रदर्शित अध्ययन कार्य तीन चीजों पर आधारित था जैसे मधुमेह की दवा (medication of diabetes), असामान्य प्रयोगशाला परीक्षण (Abnormal Laboratory Tests) और मधुमेह निदान (Diagnosis of Diabetes)। उन्होंने अपने एल्गोरिथ्म की तुलना, विभिन्न मशीन लर्निंग एल्गोरिथ्म (जैसे की सपोर्ट वेक्टर मशीन (Support Vector Machine-SVM) (एसवीएम), के-नियरेस्ट नेबर (K-Nearest Neighbour-KNN) (केएनएन) एल्गोरिथ्म, डिसिशन ट्री (Decision Tree), नैव बैस क्लासिफायर (Naïve Bayes Classifier), लोजिस्टिक रिग्रेशन (Logistic Regression) से, एक्यूरेसी स्पेसिफिसिटी और सेंसिटिविटी के आधार पर की थी।

शोधकर्ताओं ने मधुमेह के वर्गीकरण (Classification) और प्रीडिक्शन के लिए न्यूरल नेटवर्क (Neural Network) का भी अध्ययन किया, शोधकर्ताओं ने पाया कि, ट्रेनिंग डाटासेट से उत्पन्न समानताओं के आधार पर न्यूरल नेटवर्क में मधुमेह की सटीक प्रीडिक्शन/भविष्यवाणी करने की क्षमता है। दीजाना सेजिनोवी एट अल (Dijana Sejdinovi et al.) [6] ने टाइप-2 डायबिटीज मेल्लिटस (T2DM) के वर्गीकरण के लिए न्यूरल नेटवर्क्स का उपयोग किया था। तत्पश्चात मधुमेह को प्रेडिक्ट करने के लिए पारास्तो रहिमलो एट अल (Parastoo Rahimloo et al.) [7] ने आर्टिफिशियल न्यूरल नेटवर्क और लॉजिस्टिक रिग्रेशन का उपयोग करके एक हाइब्रिड न्यूरल नेटवर्क मॉडल प्रस्तुत किया। महमूद हैदरी एट अल (Mahmoud Heydari et al.) [8] [Classification] ने डायबिटीज डाटासेट पर विभिन्न क्लासिफिकेशन/वर्गीकरण के एल्गोरिथ्म का अध्ययन किया और आर्टिफिशियल न्यूरल नेटवर्क्स के उपयोग से उन्होंने आशाजनक परिणाम

पाए, पर और बेहतर परिणामों के लिए उनके शोध अध्ययन में हाइपर पैरामीटर्स को ऑप्टिमाइज करने की गुंजाइश थी। पिछले शोध कार्यों की समीक्षा से निष्कर्ष निकलता है कि पारंपरिक टेक्निक्स [9] द्वारा हासिल की गई सटीकता पर सुधार करने की गुंजाइश है। और न्यूरल नेटवर्क्स, पारम्परिक मशीन लर्निंग क्लासिफायर्स [10] की तुलना में सटीक परिणाम देने में सक्षम है।

पिछले शोध कार्यों के अध्ययन से पता चलता है कि क्लासिफायर्स की सटीकता और परफॉरमेंस को बेहतर बनाने के लिए, शोधकर्ताओं ने अपने द्वारा बनाये गए प्रेडिक्शन मॉडल में विभिन्न कारकों को शामिल किया था, जिस से उनके मॉडल की जटिलता बढ़ गयी, फलस्वरूप उनके मॉडल का कम्प्यूटेशन टाइम भी बढ़ गया।

इस समस्या को हल करने के लिए प्रस्तुत शोध पत्र में हमने प्रयास किया है कि मॉडल की जटिलता को बढ़ाये बिना, क्लासिफिकेशन परफॉरमेंस और कम्प्यूटेशन टाइम में सामंजस्य बिठाया जाए, साथ ही हमने एल्गोरिथ्म की एक्यूरेसी का तुलनात्मक अध्ययन भी किया है। आँकी गयी समस्या के निवारण के लिए हमने विभिन्न डाटा इम्प्यूटेशन टेक्निक्स का प्रयोग किया है, जैसे कि अनुपस्थित आंकड़ों एवं डाटा ऑउटलायर्स का निरीक्षण तथा अनुपस्थित आंकड़ों को संजोना। अनुपस्थित डाटा को किस प्रकार संजोना है, यह समस्या एक शोधकर्ता के लिए जटिल समस्या होती है, और एक मजबूत डेटा वर्गीकरण मॉडल बनाने के लिए, इन अनुपस्थित आंकड़ों को संजोना अति आवश्यक भी है। इस लिए, हमारा शोध सबसे पहले अनुपस्थित आंकड़ों एवं ऑउटलायर्स के निरीक्षण के लिए है [3], निरीक्षण के उपरांत उन अनुपस्थित आंकड़ों को सही करने के पश्चात, हमने विभिन्न वर्गीकरण तकनीकों का नए डेटासेट पर प्रयोग किया। एएनएन (ANN) को इस्तेमाल करके, मधुमेह (डायबिटीज) डिटेक्शन के सटीकता स्तर को ज्ञात किया और फिर एएनएन

(ANN) के सटीकता स्तर की पारंपरिक क्लासिफायर के सटीकता स्तर से तुलना की गयी।

डेटासेट में दिए गए मानों में से असामान्य मानों को ऑउटलायर (Outlier) कहा जाता है, वे या तो बहुत ज्यादा या बहुत कम मान रखते हैं, और अपेक्षित सीमा में नहीं आते हैं। यह ऑउटलायर (Outlier) किसी भी माप में परिवर्तनशीलता या प्रयोगात्मक त्रुटि का परिणाम हो सकते हैं। इस लिए एक शोधकर्ता को ऑउटलायरस को बड़ी सावधानी से संभालना होता है, विशेष रूप से चिकित्सा डेटा में, जहां यह बीमारी की भविष्यवाणी की सटीकता के लिए बहुत महत्वपूर्ण है। सामान्य तौर पर शोधकर्ता अनुपस्थित आंकड़ों [9] [11] को नियंत्रित करने के लिए माध्य या माध्यिका की गणना करते हैं, और इस दृष्टिकोण को एक अन्य शोधकर्ता कांग एट अल (Kang et al.) [12] ने सुधारा, जो कि एनेस्थेसियोलॉजी और दर्द चिकित्सा विभाग में विभिन्न प्रकार के अनुपस्थित आंकड़ों का अध्ययन करने वाले एक शोधकर्ता है।

प्रस्तुत कार्य का उद्देश्य है कि डाटा सेट में डाटा से जुड़ी त्रुटियों को हटाकर एक सुदृढ़ डाटासेट का निर्माण करना, तत्पश्चात डायबिटीज के डायग्नोसिस की एक्यूरेसी को ज्ञात करने के लिए डाटासेट पर मशीन लर्निंग के विभिन्न क्लासिफिकेशन एल्गोरिथ्म का प्रयोग करना और अंततः डायबिटीज डिटेक्शन में क्लासिफिकेशन एल्गोरिथ्म की एक्यूरेसी को आर्टिफिशियल न्यूरल नेटवर्क की एक्यूरेसी से तुलना करना है।

3. वर्गीकरण तकनीकें :

प्रस्तुत अध्ययन में वर्गीकरण तकनीकों की सटीकता जांचने के लिए हमने UCI रिपॉजिटरी द्वारा जारी किये गए पीमा इंडियन (PIMA INDIAN) नामक डाटा सेट [1] पर डाटा इंफ्यूटेशन तकनीक (क्लास-वाइज मीन, तथा 2% विनसोराइजेशन) का उपयोग किया है। उसके बाद तुलनात्मक अध्ययन के लिए इस डाटा सेट पर हमने मशीन लर्निंग की विभिन्न वर्गीकरण तकनीकों का उपयोग किया और

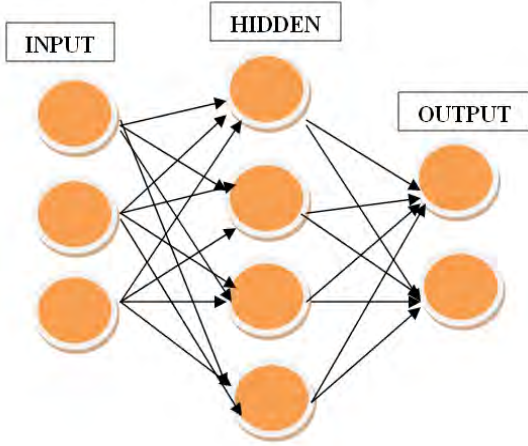
मधुमेह निदान की सटीकता जाँची है। प्रस्तुत शोध में प्रयोग की गयी तकनीक इस प्रकार है, सपोर्ट वेक्टर मशीन (Support Vector Machine), के-नियरेस्ट नेबर (K-Nearest Neighbour) एल्गोरिथ्म, डिसिशन ट्री (Decision Tree), आर्टिफिशियल न्यूरल नेटवर्क (Artificial Neural Network), नैव बैस क्लासिफायर (Naïve Bayes Classifier), इस्तेमाल की गयी इन तकनीकों की चर्चा नीचे की गयी है।

3.1. आर्टिफिशियल न्यूरल नेटवर्क (ANN)

आर्टिफिशियल न्यूरल नेटवर्क ऐसा डीप लर्निंग सिस्टम है जो वर्गीकरण (classification) और प्रतिगमन (Regression) विश्लेषण [13] के लिए उपयोग किया जाता है, इस सिस्टम की कार्यशैली मानव मस्तिष्क के समान है, जहां पर न्यूरोन्स परस्पर जुड़े होते हैं और समस्याओं को हल करने के लिए सामूहिक रूप से काम करते हैं।

ANN में तीन लेयर होती हैं, इनपुट (Input), हिडन (Hidden) और आउटपुट (Output) लेयर होते हैं।

1. इनपुट (Input) लेयर (Layer): डेटा को प्राप्त करती है और उस डाटा को नेटवर्क के अगले भाग को सौंप देती है।
2. हिडन (Hidden) लेयर (Layer): इनपुट डाटा तथा इनपुट लेयर न्यूरोन्स और हिडन लेयर न्यूरोन्स के बीच के कनेक्शन वेट्स, हिडन लेयर का कार्य निर्धारण करते हैं। इनपुट लेयर और हिडन लेयर के बीच के कनेक्शन वेट्स, निर्धारित करते हैं कि कब हिडन लेयर न्यूरोन सक्रिय (एक्टिव) होगा और कब वह निष्क्रिय (इनएक्टिव) होगा।
3. आउटपुट (Output) लेयर (Layer) हिडन लेयर के न्यूरोन्स का आउटपुट तथा हिडन लेयर और आउटपुट लेयर के बीच के कनेक्शन वेट्स, आउटपुट लेयर के न्यूरोन्स का कार्य निर्धारण करते हैं।



चित्र 1. आर्टिफिशियल न्यूरल नेटवर्क (ANN) की संरचना।

किसी भी फीड फोरवर्ड न्यूरल नेटवर्क में हर लेयर अपने से पिछली लेयर के आउटपुट पर निर्भर करती है और सभी लेयर्स के वेट्स को समायोजित करके मशीन कुछ सीख पाती है।

इस संरचना में पहली परत यानी इनपुट लेयर को इस प्रकार लिखा जा सकता है:

$$h1 = g1 (W1 * x + b1)$$

जहाँ, $g1$ एक्टिवेशन फंक्शन है, $W1$ का अर्थ है वेट्स (Weights) और $b1$ बायस टर्म है

इसी प्रकार, हिडन लेयर (दूसरी लेयर) को इस प्रकार लिखा जा सकता है:

$$h2 = g2 (W2 * h1 + b2)$$

और थर्ड लेयर (आउटपुट लेयर) को इस प्रकार लिखा जा सकता है:

$$y = g3 (W3 * h2 + b3)$$

अब, $h1$ और $h2$ का प्रयोग करके हम लिख सकते हैं

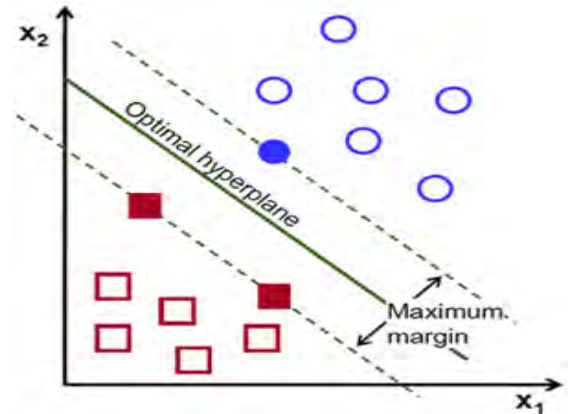
$$y = g3 (W3 * g2 (W2 * g1 (W1 * x)) + [b + \dots])$$

ऊपर दिए गए समीकरणों के आधार पर सभी लेयर्स के बीच त्रुटि (Error) का आंकलन किया

जाता है और इसके आधार पर सभी हिडन नोड्स के कनेक्शन वेट्स पर करेक्शन (Correction) लगायी जाती है, और इस प्रकार इनपुट नोड्स प्रत्येक अवलोकन (ऑब्जरवेशन) के लिए बार-बार सेट किये जाते हैं।

3.2. सपोर्ट वेक्टर मशीन (SVM)

एक सपोर्ट वेक्टर मशीन (SVM) विभिन्न डेटा वर्गों को वर्गीकृत करने के लिए हाइपर-प्लेन का उपयोग करती है। एसवीएम का मॉडल ट्रेनिंग डाटा पर बनाया जाता है और टेस्ट डाटा से एक हाइपर प्लेन जनरेट किया जाता है। एसवीएम मॉडल का उपयोग करके हम, डेटा मैट्रिक्स में उस स्थान का पता लगते हैं, जहां विभिन्न डेटा समूहों को व्यापक रूप से अलग किया जा सकता है, और इस स्थान को हाइपर-प्लेन निर्धारित करता है। नीचे दिए गए चित्र-2 में लाल और नीले रंग के दो ट्रेनिंग डाटा पॉइंट्स दिखाए गए हैं।



चित्र 2. सपोर्ट वेक्टर मशीन (SVM)

हाइपर-प्लेन को रैखिक रूप से परिभाषित करने के लिए हमें एक ऐसा प्लेन बनाना होता है, जो डाटा को अधिकतम मार्जिन से अलग कर सके, यह एक आदर्श हाइपर-प्लेन (Optimal hyper plane) कहलाता है। यह हाइपर प्लेन रेखिक भी हो सकता है और अरेखिक लीनियर भी।

3.3. के-नियरेस्ट नेबर (KNN)

के नियरेस्ट नेबर एल्गोरिदम सभी मशीन लर्निंग एल्गोरिदम में सबसे सरल हैं, इन एल्गोरिदम के पीछे मूल अवधारणा यह है कि, प्रशिक्षण सेट को याद रखना है और फिर अपने निकटतम पड़ोसियों के लेबल के आधार पर प्रशिक्षण सेट में हर नए उदाहरण बिंदु के लेबल का निर्धारण करना है। इस वर्गीकरण की तकनीक में डोमेन बिंदु लेबलिंग को परिभाषित करने के लिए उपयोग किए जाते हैं, क्योंकि हर डोमेन बिंदु के पास के अधिकांश बिंदुओं में एक ही लेबल होता है, इसलिए नियरेस्ट नेबर एल्गोरिदम में यह माना जाता है कि वह डेटा बिंदु जो रिफरेन्स बिंदु से न्यूनतम दूरी पर स्थित होगा, डेटा बिंदु उसी वर्ग (बर्से) का होगा।

के-नियरेस्ट नेबर एल्गोरिथ्म में, डेटा बिंदुओं के बीच की दूरी को खोजने के लिए यूक्लिडियन डिस्टेंस का उपयोग किया जाता है, इस फंक्शन का फॉर्मूला नीचे दिया गया है:

$$d(q,p) = \sqrt{(q_1-p_1)^2 + (q_2-p_2)^2 \dots + (q_n-p_n)^2}$$

$$= \sqrt{\sum_{i=1}^n (q_i-p_i)^2}$$

जहां, n आयामों की संख्या है, या हमारी मशीन सीखने की विशेषताएं हैं।

3.4. डिसिशन ट्री (Decision Tree)

डिसिशन ट्री, मशीन लर्निंग की तकनीकों में से,

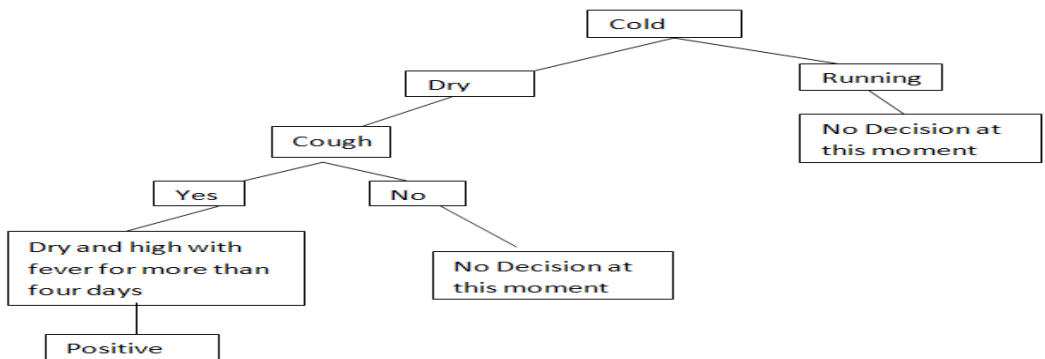
सुपरवाइसड लर्निंग की श्रेणी की एक तकनीक है। डिसिशन ट्री, एक पेड़ जैसा आकार होता है, जहां नोड एक फीचर पर टेस्ट का प्रतिनिधित्व करते हैं, क्लास लेबल, डिसिशन ट्री की लीफ नोड्स द्वारा दर्शाए जाते हैं और ब्रांच उन फीचर्स के कंजक्शन को दर्शाते हैं, जो उन क्लास लेबल्स को कनेक्ट करती हैं। डिसिशन ट्री की जड़ (रूट नोड—Root Node) से डिसिशन ट्री की पत्ती (लीफ नोड – leaf node) तक के रास्ते वर्गीकरण नियमों (Classification rules) को दर्शाते हैं। सभी फीचर्स की गणना के बाद वर्गीकरण पर निर्णय लिया जाता है। नीचे चित्र 3. में डिसिशन ट्री की संरचना के द्वारा वर्गीकरण करने के लिए एक उदाहरण दिया गया है कि व्यक्ति कोरोना वायरस से संक्रमित है या नहीं।

व्यक्ति Covid19 से संक्रमित है या नहीं (ऊपर निर्णय वृक्ष का प्रदर्शन करने के लिए सिर्फ एक उदाहरण है और चिकित्सा दिशा निर्देशों के अधीन नहीं है)

3.5. नैव बैस क्लासिफायर

नैव बैस क्लासिफायर, प्रायकिता के आधार पर वर्गीकरण करने की एक तकनीक है, जो डेटा को वर्गीकृत करने के लिए बैस प्रमेय (Theorem) का उपयोग करती है।

बैस प्रमेय (Theorem) के हिसाब से घटना ए के होने की सम्भावना, जबकि घटना बी पहले ही घटित हो चुकी हो, कुछ नीचे दिए गए फॉर्मूले से दे जाती है



चित्र 3. डिसिशन ट्री ;(Decision Tree) की संरचना

यहाँ, B प्रमाण (Evidence) है और A परिकल्पना (हाइपोथिसिस— Hypothesis) है। यहाँ पूर्वधारणा यह है कि प्रीडिक्टर्स या फीचर्स मतलब स्वतंत्र हैं। कि एक प्रीडिक्टर्स या फीचर्स की उपस्थिति दूसरे को प्रभावित नहीं करती है। इसलिए इसे नैव (Naïve) कहा जाता है।

डेटासेट :

इस शोध पत्र में, इस्तेमाल किया गया डेटासेट PIMA भारतीय डेटासेट है। ऐसा डेटासेट यूसीआई रिपॉजिटरी [1] में सभी के लिए उपलब्ध है। इस डेटासेट के चुनाव के पीछे मुख्य कारण यह है इसमें गर्भावस्था और रक्त प्लाज्मा विवरणों के अलावा

$$P(A|B) = \frac{P(B|A)P(A)}{P(B)}$$

बुनियादी, मधुमेह से जुड़े आंकड़े उपस्थित हैं, जैसे कि उम्र, अधिक वजन, बीएमआई, सीरम, इंसुलिन, डायस्टोलिक रक्तचाप, और इसके अलावा इसमें डायबिटीज पेडिग्री फंक्शन (DPF) भी शामिल है, जो रिश्तेदारों में डायबिटीज मेल्लिटस के इतिहास तथा रोगी का उन रिश्तेदारों के आनुवंशिक संबंधों से संबंधित डेटा है।

इस डेटासेट में भी कहीं कहीं पर कुछ आंकड़ों की अनुपस्थिति थी (Missing data values), जो कि मजबूत तकनीकी वर्गीकरण मॉडल तैयार करने के लिए शोधकर्ताओं द्वारा सबसे बड़ी चुनौती है। इसलिए, अपनी शोध में हमने सबसे पहले अनुपस्थित आंकड़ों को और आउटलायर्स को संभालने के लिए प्रयास किये हैं। इस पत्र में, हमने, अनुपस्थित डाटा के आंकलन के लिए क्लास-वाइज मीन डाटा इम्प्यूटेशन तकनीक का उपयोग किया है और आउटलायर को संभालने के लिए 2% विनसोराइजेशन तकनीक का उपयोग किया है। हमने अपने पिछले अध्ययन [3] में पाया था कि क्लास-वाइज मीन

तकनीक की सटीकता (accuracy), संवेदनशीलता (sensitivity) और विशिष्टता (Specificity) के आधार पर सर्वोत्तम है। अंततः डाटा सेट को क्लास-वाइज मीन और विनसोराइजेशन तकनीक से सुनिश्चित करने के पश्चात (निरिक्षण के उपरांत), हमने विभिन्न वर्गीकरण तकनीकों का डेटासेट पर उपयोग किया और उन वर्गीकरण तकनीकों की मधुमेह (डायबिटीज) डिटेक्शन में सटीकता (एक्यूरेसी) का तुलनात्मक अध्ययन किया है।

4.1 परफॉर्मेंस इवैल्यूएशन

प्रस्तुत कार्य में, क्लासिफिकेशन तकनीक की एक्यूरेसी को उनके परफॉर्मेंस फैक्टर के रूप में माना गया है। क्लासिफिकेशन तकनीक सटीकता को नापने के लिए हमे वास्तविक सकारात्मक (ट्रू पॉजिटिव – TP) और वास्तविक नकारात्मक (ट्रू नेगेटिव – TN) वर्गों का उपयोग, मरीज के वर्गीकरण के लिए करना पड़ता है, कि वह व्यक्ति मधुमेह से पीड़ित है या नहीं। दूसरे शब्दों में क्लासिफिकेशन तकनीक/मैथड की एक्यूरेसी/सटीकता निकलने के लिए, मामलों को सही ढंग से वर्गीकृत करने की दर को निकाला जाता है। एक्यूरेसी (Accuracy) निकलने का फार्मूला नीचे दिया गया है

$$\text{accuracy} = \frac{TP+TN}{TP+TN+FP+FN} \quad (1)$$

$$\text{sensitivity} = \frac{TP}{TP+FN} \quad (2)$$

$$\text{specificity} = \frac{TN}{TN+FP} \quad (3)$$

$$\text{precision} = \frac{TP}{TP+FP} \quad (4)$$

$$\text{recall} = \frac{TP}{TP+FN} \quad (5)$$

$$F - \text{measure} = \frac{2 * \text{precision} * \text{recall}}{\text{precision} + \text{recall}} \quad (6)$$

यहाँ,

TP: अर्थात वास्तविक सकारात्मक (ट्रू पॉजिटिव) जिसका अर्थ है सही ढंग से वर्गीकृत सकारात्मक मामले।

TN: अर्थात वास्तविक नकारात्मक (ट्रू नेगेटिव) जिसका अर्थ है सही ढंग से वर्गीकृत नकारात्मक मामले।

FP: अर्थात काल्पनिक सकारात्मक (फाल्स पॉजिटिव) जिसका अर्थ है गलत तरीके से वर्गीकृत नकारात्मक मामले।

FN: अर्थात काल्पनिक नकारात्मक (फाल्स नेगेटिव) जिसका अर्थ है गलत तरीके से वर्गीकृत सकारात्मक मामले।

हमने डाटा मॉडल को के- फोल्ड क्रॉस (K-Fold Cross) वेलिडेशन/मान्यीकरण (Validation) तकनीक का उपयोग करके मान्य किया है, जो डेटा मॉडल के प्रदर्शन को मान्य करने के लिए सबसे अधिक बार उपयोग की जाने वाली विधि है। इस शोध कार्य में, 10- फोल्ड क्रॉस (10-Fold Cross)- वेलिडेशन का उपयोग किया गया है, जिसमें प्रारंभिक डेटासेट को 10 उप-नमूनों में विभाजित किया जाता है, और एक भाग का उपयोग परीक्षण के उद्देश्य के लिए किया जाता है तथा शेष नौ भागों को प्रशिक्षण के लिए उपयोग किया जाता है। इस प्रक्रिया को दस बार दोहराया गया और औसत सटीकता को मॉडल की अंतिम सटीकता के रूप में लिया गया।

मैथोडोलॉजी :

इस पत्र में, हमने, अनुपस्थित डाटा के आंकलन के लिए क्लास-वाइज मीन डाटा इम्प्यूटेशन तकनीक का उपयोग किया है और ऑउटलायर को सँभालने के लिए 2% विनसोराइजेशन तकनीक का उपयोग किया है। हमने अपने पिछले अध्ययन [3] में पाया था की क्लास-वाइज मीन तकनीक सटीकता (accuracy), संवेदनशीलता (sensitivity)

और विशिष्टता (Specificity) के आधार पर सर्वोत्तम है। अंततः डाटा सेट को क्लास-वाइज मीन और विनसोराइजेशन तकनीक से सुनिश्चित करने के पश्चात (निरिक्षण के उपरांत), हमने विभिन्न वर्गीकरण तकनीकों का डेटासेट पर प्रयोग किया और उन वर्गीकरण तकनीकों की मधुमेह (डायबिटीज) डिटेक्शन में सटीकता (एक्यूरेसी) का तुलनात्मक अध्ययन किया है।

प्रयोग को दो चरणों में आयोजित किया गया था, पहले चरण/स्टेज में डाटा प्रीप्रोसेसिंग की गयी है, जहां डुप्लीकेट और अनुपस्थित डाटा के लिए डेटासेट की जाँच की गई, और डेटा इंप्यूटेशन (Data Imputation) के लिए, क्लास वाइज मीन (Class wise mean) तकनीक [3] का उपयोग किया गया। उसके बाद डेटा को दो भागों में विभाजित किया गया, पहले विभाजन में मिनि-मैक्स स्केलर (Min-Max Scalar) को फीचर डाटा सेट पर अप्लाई कर के उनकी वैल्यूज को 0 से 1 के बीच सेट कर दिया गया, और दूसरी बार टेस्टिंग तथा ट्रेनिंग के लिए फीचर डाटा सेट को विभाजित किया गया।

दूसरे चरण में तीन परत के एएनएन (ANN) मॉडल का निर्माण किया गया था, जहां इनपुट (input) और अदृश्य (hidden) परत में न्यूरॉन्स और एक्टिवेशन फंक्शन थे और आउटपुट (output) परत/लेयर्स में सिग्मोइड (Sigmoid) फंक्शन का इस्तेमाल किया है। उसके बाद एएनएन (ANN) के आउटपुट को के-फोल्ड क्रॉस (K-fold cross) वेलिडेशन तकनीक से वैलिडेट किया गया। फिर पायथन साइकेट लाइब्रेरी का उपयोग करके, एएनएन (ANN) की सटीकता की तुलना अन्य क्लासिफायर (यानी SVM, KNN, डिसीजन ट्री और नैव बैस क्लासिफायर) की सटीकता के साथ की गयी है।

चरण 1 का सूडो कोड (Pseudo Code): डाटा प्रीप्रोसेसिंग स्टेज और चरण 2: एएनएन बनाने के चरणों का सूडो कोड (Pseudo Code) नीचे प्रस्तुत किया गया है:

चरण/स्टेज-1 :	चरण/स्टेज-2 :
डाटा प्रीप्रोसेसिंग स्टेज का सूडो कोड (Pseudo Code): इनपुट: बाइनरी लॉस फंक्शन और स्टोकोस्टिक ग्रेडिएंट ऑप्टिमाइजर इनपुट: पिमा इंडियन डेटासेट .csv के रूप में डेटासेट की प्री-प्रोसेसिंग (Pre Processing): - डुप्लिकेट के लिए जाँच कि यदि मौजूद है तो हटा दिया गया है - अनुपस्थित आंकड़ों को संभालने के लिए क्लास वाइज डाटा इम्यूटेशन तकनीक का प्रयोग किया गया। - ऐरे/सरणी (Array) में उपस्थित डाटा से न्यूरल नेटवर्क का निर्माण - स्प्लिट डेटासेट डी(D): ☐ — डेटासेट में x और ल, जहां x स्वतंत्र फीचर है और ल निर्भर फीचर है। — मिन-मैक्स स्केलर विधि का प्रयोग:फीचरडेटासेट में 0 और 1 के बीच रखने के लिए। - दोबारा/फिर से डेटासेट(D) स्प्लिट: — प्रशिक्षण डेटा 80% — परीक्षण डेटा 20%	ANN मॉडल बनाने का सूडो कोड (Pseudo Code): इस मॉडल में इनपुट, आउटपुट और हिडन लेयर के रूप में तीन लेयर है। - इनपुट परत/लेयर में 15 न्यूरॉन्स और एक्टिवेशन फंक्शन हैं, - दूसरी परत/लेयर में 17 न्यूरॉन्स और एक्टिवेशन फंक्शन हैं और - आउटपुट एक्टिवेशन में 1 न्यूरॉन और सिग्मॉइड फंक्शन होते हैं। एएनएन(ANN) के आउटपुट को के-फोल्ड क्रॉस (K-fold cross) वैलिडेशन तकनीक से वैलिडेट किया गया। मॉडल में प्रशिक्षण डाटा पर बाइनरी लॉस फंक्शन और स्टोकोस्टिक ग्रेडिएंट ऑप्टिमाइजर का उपयोग करके,उसकी सटीकता/एक्यूरेसी को मापने के लिए संकलित किया गया था और सटीकता को मैट्रिक्स के रूप में संजोया गया था।

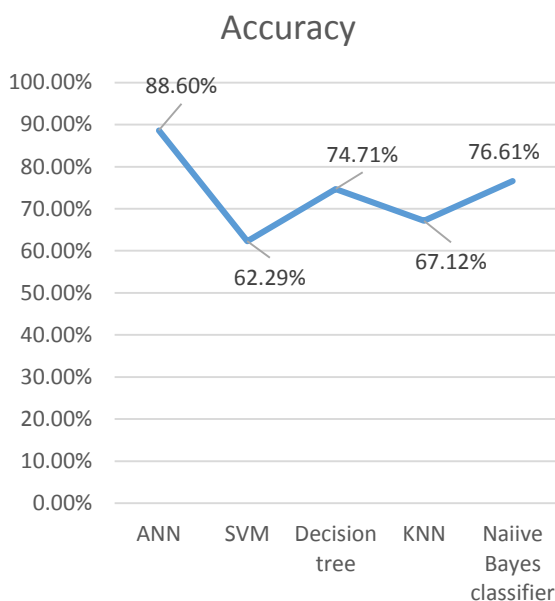
6. परिणाम और चर्चा

प्रत्येक क्लासिफायर की सटीकता/एक्यूरेसी तालिका-2 में नीचे दी गई है, परिणामों से यह स्पष्ट है कि आर्टिफिशियल न्यूरल नेटवर्क (एएनएन), चुने हुए डेटासेट पर 88.6% की सटीकता/एक्यूरेसी दर्शाते हुए सबसे अच्छा प्रदर्शन किया है

तालिका 2. : विभिन्न वर्गीकरण एल्गोरिदम की सटीकता (%)

वर्गीकरण तकनीक	सटीकता
आर्टिफिशियल न्यूरल नेटवर्क	88.6%
सपोर्ट वेक्टर मशीन	62.29%
डिसिशन ट्री	74.71%
के-नियेरेस्ट नेबर	67.12%
नैव बैस क्लासिफायर	76.61%

चूंकि आर्टिफिशियल न्यूरल नेटवर्क (ANN) केवल 0 और 1 के बीच संभाव्यता संख्या (Probability असंनम) को मान देता है, इसलिए 0.5 की सीमा को डेटा को 1 और 0. के रूप में वर्गीकृत करने के लिए सेट किया गया था। यदि मान 0.5 से ऊपर था तो इसे 1 के रूप में वर्गीकृत किया गया था अन्यथा 0 के रूप में वर्गीकृत किया गया था आर्टिफिशियल न्यूरल नेटवर्क (ANN) पारंपरिक मॉडल की तुलना पारंपरिक मशीन लर्निंग क्लासिफायर के साथ की गई थी।



चित्र 4. एनएन के साथ तुलना में विभिन्न क्लासिफायर का सटीकता

वास्तव में, बहुत मुश्किल है यह कह पाना की कब कौनसी वर्गीकरण तकनीक सबसे अच्छा प्रदर्शन करेगी, लेकिन विशेषज्ञों ने देखा कि डीप लर्निंग तकनीक में सर्वोत्तम संभव परिणाम प्राप्त करने की क्षमता है। भविष्य के कार्यों में, रक्त शर्करा के स्तर की भविष्यवाणी करने के लिए समय श्रृंखला (time series) डेटा को उपयोग करके, डीप लर्निंग मॉडल को विकसित किया जा सकता है। जो इंसुलिन के स्तर को बनाए रखने में महत्वपूर्ण योगदान हो सकता

है। डीप लर्निंग की क्षमता इतनी जबरदस्त है कि यह भविष्य में एक बड़ी प्रगति ले सकती है।

7. निष्कर्ष :

इस पत्र में, पारंपरिक वर्गीकरण तकनीकों की तुलना आर्टिफिशियल न्यूरल नेटवर्क (ANN) के साथ की गयी है। तालिका -2 में दिए गए परिणामों के आधार पर, यह निष्कर्ष निकाला गया है कि आर्टिफिशियल न्यूरल नेटवर्क (ANN) में मधुमेह के वर्गीकरण के लिए उच्च स्तर की सटीकता (88.6%) पाई गई, जिसमें Pima Indian Dataset की सभी विशेषताएं/फीचर्स शामिल हैं। अन्य शोधों में, शोधकर्ताओं ने अधिक सटीकता पाई, लेकिन वे आमतौर पर उन्होंने अनुपस्थित आंकड़े/डेटा या आउटलायर्स को नजरअंदाज कर दिया था। प्रस्तुत शोध कार्य में क्लास वाइज मीन (class wise mean) [3] का उपयोग करके अनुपस्थित आंकड़े/डेटा और आउटलायर्स को संभालते हुए, और मॉडल की जटिलता को बिना बढ़ाये, हमने पाया कि मधुमेह के वर्गीकरण के लिए आर्टिफिशियल न्यूरल नेटवर्क (ANN) की तकनीक, सर्वोत्तम है और सटीक रिजल्ट्स देती है।

8. प्रमुख शब्दों की तालिका :

Technical Terms (English)	तकनीकी शब्द (हिन्दी)
Classification Algorithm	वर्गीकरण तकनीक
Support Vector Machine - SVM	सपोर्ट वेक्टर मशीन
K&Nearest Neighbor - KNN	के-नियरेस्ट नेबर
Decision Tree	डिसिशन ट्री
Artificial Neural Network - ANN	आर्टिफिशियल न्यूरल नेटवर्क्स
Naïve Bayes Classifier	नैव बैस क्लासिफायर
Logistic Regression	लोजिस्टिक रिग्रेशन

Missing data values	आंकड़ों की अनुपस्थिति
Sequential Covering Approach	सिक्वेन्शियल कवरिंग अप्रोच
Medication of Diabetes	मधुमेह की दवा
Abnormal Laboratory Tests	असामान्य प्रयोगशाला परीक्षण
Diagnosis of Diabetes	मधुमेह निदान
Classification	वर्गीकरण
Accuracy	सटीकता
Sensitivity	संवेदनशीलता
Specificity	विशिष्टता
Probability	सम्भावना
Data Imputation	डेटा इंप्यूटेशन
Time series	समय श्रृंखला

References:

- [1] UCI repository, Pima Indian Dataset.
- [2] Wenqian, Shuyu Chen, Hancui Zhang, and Tianshu Wu Chen, "A hybrid prediction model for type 2 diabetes using K-means and decision tree," in Software Engineering and Service Science (ICSESS), pp. 386-390, 2017.
- [3] Sofia Goel and Sudhansh Sharma, "Advanced Data Imputation Techniques for Predicting Type 2 Diabetes using Machine Learning," "International Journal of Innovative Technology and Exploring Engineering (IJITEE), ISSN: 2278-3075, vol.9, issue 2, December 2019.
- [5] Nahla, Andrew P. Bradley, and Mohamed Nabil H. Barakat Barakat, "Intelligible support vector machines for diagnosis of diabetes mellitus," IEEE transactions on information technology in biomedicine, vol. 14, no. 4, pp. 1114-1120, July 2010.
- [5] Xie W, Xu L, He X, Zhang Y, You M, Yang G, Chen Y Zheng T, "A machine learning-based framework to identify type 2 diabetes through electronic health records," International journal of medical informatics, vol. 97, pp. 120-127, Jan 2017.
- [6] Dijana Sejdinovic et al., "Classification Of Prediabetes And Type 2 Diabetes Using Artificial Neural Network", CMBEIHS Springer, Singapore. pp. 685-689, March 2017.
- [7] Rahimloo, Parastoo, and Ahmad Jafarian, "Prediction of Diabetes by Using Artificial Neural Network, Logistic Regression Statistical Model and Combination of Them", Bulletin de la Société Royale des Sciences de Liège, 85, pp. 1148-1164, Jan 2016.
- [8] Mahmoud, Mehdi Teimouri, Zainab Alhoda Heshmati, and Seyed Mohammad Alavinia Heydari, "Comparison of various classification algorithms in the diagnosis of type 2 diabetes in Iran," International Journal of Diabetes in Developing Countries, vol. 36, pp. 167-173, 2016.
- [9] Longfei, Senlin Luo, Jianmin Yu, Limin Pan, and Songjing Chen Han, "Rule extraction from support vector machines using ensemble learning approach: an application for diagnosis of diabetes," IEEE journal of biomedical and health informatics, vol. 19, pp. 728-734., 2015.
- [10] Ali, Tinna B. Aradóttir, Alexander R. Johansen, Henrik Bengtsson, Marco Fraccaro, and Morten Mørup Mohebbi, "A deep learning approach to adherence detection for type 2 diabetics," in Engineering in Medicine and Biology Society (EMBC), vol. 39, pp. 2896-2899, 2017.
- [11] Xie W, Xu L, He X, Zhang Y, You M, Yang G, Chen Y Zheng T, "A machine learning-based framework to identify type 2 diabetes through electronic health records," International journal of medical informatics, vol. 97, pp. 120-127, Jan 2017.
- [12] Hyun Kang, "The prevention and handling of the missing data.," Korean journal of anesthesiology, pp. 402-406, 2013.
- [13] Aliza Ahmad, Aida Mustapha, Eliza Dianna Zahadi, Norhayati Masah, and Nur Yasmin Yahaya, "Comparison between Neural Networks against Decision Tree in Improving Prediction Accuracy for Diabetes Mellitus," in International Conference on Digital Information Processing and Communications, pp. 537-545, 2011.

नैचुरल लैंग्वेज प्रोसेसिंग में ऊर्दू स्टेमर का विकास

Development of Urdu Stemmer in Natural Language Processing

वैशाली गुप्ता¹, निशीथ जोशी², इति माथुर³

¹ आई.पी.एस. एकैडमी, आई.ई.एस., इंदौर, मध्य प्रदेश

^{2,3} आपाजी संस्थान, वनस्थली विद्यापीठ, राजस्थान

¹ vaishali.gupta77@gmail.com, ² nisheeth.joshi@rediffmail.com, ³ mathur_iti@rediffmail.com

सारांश

स्टेमिंग (Stemming) की प्रक्रिया मार्फोलोजिकल एनालाइजर (Morphological Analyser), इन्फोर्मेशन रिट्रीवल (Information Retrieval), नैचुरल लैंग्वेज प्रोसेसिंग (Natural Language Processing) एवं स्पेल चेकिंग (Spell Checking) इत्यादि में बहुत महत्वपूर्ण स्थान रखती है। स्टेमिंग के द्वारा हम किसी भी शब्द के “स्टेम” (आधार अथवा मूल) को प्रथक कर सकते हैं। दूसरे शब्दों में हम कह सकते हैं कि प्रेफिक्स (prefix) (उपसर्ग) एवं /अथवा सफिक्स (suffix) (प्रत्यय) को किसी शब्द से हटा कर ‘मूल’ शब्द प्राप्त करना ही स्टेमिंग है। स्टेमिंग से प्राप्त मूल शब्द इन्फ्लेक्शनल (Inflectional) अथवा डेरिवेशनल मार्फोलोजी (Derivational Morphology) को प्रदर्शित करता है। मार्फोलोजी से अर्थ है किसी भी शब्द की सम्पूर्ण रचना। इन्फ्लेक्शनल मार्फोलोजी में मूल शब्द की श्रेणी उसके वास्तविक शब्द की श्रेणी के समान होती है। डेरिवेशनल मार्फोलोजी में मूल शब्द की श्रेणी उसके वास्तविक शब्द की श्रेणी से भिन्न होती है। इस शोध पत्र में, हमने नियमबद्ध इन्फ्लेक्शनल एवं डेरिवेशनल ऊर्दू स्टेमर के विकास पर दृष्टि डाली है। यहाँ हमने लॉंगेस्ट सफिक्स स्ट्रिपिंग एल्गोरिथम (Longest Suffix Stripping) का इस्तेमाल किया है जिसके द्वारा किसी भी शब्द एवं वाक्यों में से ‘स्टेम’ शब्द का प्रथकीकरण किया जा सकता है।

Abstract:

The process of ‘Stemming’ plays an important role in the field of Natural Language Processing, Morphological Analyzer, Spell Checker and Information Retrieval. Through the Stemming, we can extract ‘root’ or ‘stem’ part from the given word. In other words, we can say that Stemming is a process of removing suffix and prefix from ‘root’ or ‘stem’ word. These obtained ‘root’ words display inflectional or derivational morphology. Morphology means the complete structure of any word. In Inflectional morphology, category of ‘root’ word is similar to its actual word and category of ‘root’ word is different to its actual or given word in derivational morphology. In this research paper, we focus on the development of inflectional and derivational rule based Urdu Stemmer. Here we are using Suffix Stripping algorithm for the extraction of ‘Stem’ or ‘Root’ word from given word.)

सूचक-शब्द: स्टेम, इन्फ्लेक्शनल, डेरिवेशनल, मार्फोलोजी, अफिक्स, ऊर्दू।

Keywords: stem, inflectional, derivational, morphology, affix, Urdu

1. परिचय

आज के युग में कम्प्यूटर का महत्वपूर्ण योगदान है। हर तरह के कार्य कम्प्यूटर के द्वारा आसानी से किये जा सकते हैं। अब हम यदि भाषा की बात करें तो हम पार-बहुभाषी (cross-lingual) विश्व में रह रहे

हैं जहाँ हर व्यक्ति को हर भाषा का ज्ञान हो यह आवश्यक नहीं है। इस भाषा रूपी समस्या के समाधान के लिए शोधकर्ताओं ने "मशीन अनुवाद" नामक तकनीक को प्रस्तावित किया। इस मशीन अनुवादक सॉफ्टवेयर के द्वारा एक भाषा को दूसरी भाषा में परिवर्तित कर सकते हैं, जिससे हम कहीं भी किसी भी देश के बहुभाषाबद्ध (multilingual) लोगों से जुड़े रह सकते हैं।

इस मशीन अनुवादक के विकास में बहुत सारी समस्याएं उभर कर आईं। सर्वप्रथम समस्या शब्दों के रूपात्मक विश्लेषण की आई। रूपात्मक विश्लेषण के द्वारा किसी भी शब्द की व्याकरण संबंधी जानकारी प्राप्त की जा सकती है। इसकी दो विधियाँ हैं : इन्प्लेक्शनल एवं डेरिवेशनल। इन्प्लेक्शनल रूप में स्टेम शब्द की श्रेणी उसके वास्तविक शब्द की श्रेणी के समान होती है एवं डेरिवेशनल रूप में स्टेम शब्द की श्रेणी उसके वास्तविक शब्द की श्रेणी से भिन्न होती है। इन दोनों विधियों की प्रक्रिया को पूर्ण करने के लिए स्टेमर का विकास किया गया है।

स्टेमिंग एक तरह की प्रक्रिया है जिसमें अफिक्स (प्रत्यय) अपने स्टेम शब्द से अलग हो जाते हैं। अफिक्स के अंतर्गत ही प्रेफिक्स एवम् सफिक्स आते हैं। प्रेफिक्स वो होते हैं जो शब्द की शुरुआत में जुड़े होते हैं और सफिक्स शब्द के अंत में जुड़ा भाग होता है। जबकि स्टेम अपने आप में शब्द का आधार है। उदाहरण के लिए— प्रसन्नता, प्रसन्नमयी, प्रसन्नचित एवं अप्रसन्न, इन शब्दों के अफिक्स 'ता', 'मयी', 'चित' एवं 'अ' होंगे। इन सभी शब्दों का आधार यानि स्टेम 'प्रसन्न' होगा। यह उभयनिष्ठ आधार, विभिन्न मार्फोलोजिकल शब्दों से सम्बंधित है। इसी प्रकार ऊर्दू भाषा में — **صبر (बेसब्र)** : **بے صبری (बेसब्री)** एवं **سبردار (सब्रदार)** शब्दों के अफिक्स 'بے (बे)' एवं 'ی (ई)' एवं 'دار (दार)' होंगे और इनका स्टेम 'صبر (सब्र)' होगा।

स्टेमर वास्तव में एक एल्गोरिथ्म (कलन विधि) है, जिसके माध्यम से हम शब्दों के स्टेम को अलग

कर सकते हैं। यह नैचुरल लैंग्वेज प्रोसेसिंग में महत्वपूर्ण स्थान रखती है। यह पद्धति स्पेल चेकिंग, मशीन अनुवाद के मूल्यांकन और इन्फोर्मेशन रिट्रीवल में विशेषतः इस्तेमाल की जाती है। यहाँ हम ऊर्दू भाषा के लिए एक नियमबद्ध स्टेमर प्रस्तुत करने जा रहे हैं। ऊर्दू भाषा में कई तरह की चुनौतियों का सामना करना पड़ता है। ऊर्दू दुर्बल इन्प्लेक्शनल भाषा है और इसकी जटिल एवं समृद्ध मार्फोलोजी एवं इसकी विविधतापूर्ण प्रकृति के कारण, स्टेमर का ऊर्दू में विकास चुनौतीपूर्ण है। यहां पर इन्प्लेक्शन को हम शब्द निर्माण की प्रक्रिया के तौर पर समझ सकते हैं, जिसके अंतर्गत हम एक शब्द को विभिन्न व्याकरण की श्रेणियों में संशोधित कर सकते हैं किंतु ऊर्दू भाषा को हम दुर्बल इन्प्लेक्शन इसलिये कहते हैं क्योंकि ऊर्दू में प्रायः शब्द के इन्प्लेक्शन मौजूद नहीं होते हैं। उदाहरण के तौर पर — मशहूर : एकवचन एवं मशाहीर : बहुवचन। इसकी जटिल मार्फोलोजी के कारण, इन शब्दों में कोई भी इन्प्लेक्शन मौजूद नहीं है। अतः इन शब्दों में से स्टेम शब्द को परित्यक्त करना संभव नहीं है। इस प्रकार के शब्द अपवाद में गिने जाते हैं। किन्तु कुछ उदाहरण जैसे— बदमिजाज, बदमिजाजी एवं मिजाज—आसना शब्द को हम इन्प्लेक्शनल या डेरिवेशनल शब्दों में गिन सकते हैं। चूँकि इन शब्दों में से हम इनका स्टेम अलग कर सकते हैं। साथ ही यहाँ पर समृद्ध मार्फोलोजी से तात्पर्य है कि ऊर्दू भाषा में विभिन्न व्याकरण की श्रेणियाँ हैं जैसे कि संज्ञा (ऊर्दू में 'इस्म'), सर्वनाम ('जमीर') इत्यादि।

शब्दों से उनके स्टेम भाग को पृथक करने के अंतर्गत अंडर स्टेमिंग (Over Stemming) एवं ओवर स्टेमिंग (Over Stemming) की समस्याएं आती हैं। यदि हम किसी अफिक्स को उसके मूल से विरक्त कर दें जबकि उस शब्द से वास्तव में किसी भी अफिक्स को विरक्त करने की आवश्यकता ही ना रही हो तो ऐसी समस्या को अंडर स्टेमिंग में शामिल करते हैं। उदाहरण के लिए - 'پیشگی' (पेशगी) शब्द से यदि अफिक्स 'ی (ई)' हटा देते हैं तो स्टेम शब्द 'پیشگ' (पेशग) प्राप्त होता है जो की अर्थहीन है। इस

‘پیش’ (पेश) की बजाय हमें स्टेम ‘پیش’ (पेश) एवं अफिक्स ‘گی’ (गी) प्राप्त होना चाहिए था। अर्थात् बहुत ही छोटे अथवा कम शब्द को किसी पूर्ण शब्द से हटाना, अंडर स्टेमिंग के अंतर्गत आएगा। इसके विपरीत ओवर स्टेमिंग में जिन अफिक्स को नहीं हटाना चाहिए वो भी हट जाते हैं। उदाहरण के तौर पर ‘بدماش’ (बदमाश) शब्द से अगर ‘بد’ (बद) अलग हो जाये एवं स्टेम शब्द ‘ماش’ (माश) बचता है तो इस स्टेम शब्द ‘ماش’ (माश) का कोई तात्पर्य नहीं है। यद्यपि इस शब्द से हमें प्रेफिक्स हटाने की आवश्यकता ही नहीं थी। ओवर स्टेमिंग में प्रायः अधिक से अधिक लम्बाई के अफिक्स के हट जाने की समस्या उभर कर आती है।

2. साहित्य की समीक्षा

नैचुरल लैंग्वेज प्रोसेसिंग के क्षेत्र में, प्रथम स्टेमर लोविंस (Lovins) [1] के द्वारा 1968 में बनाया गया था। उन्होंने अंग्रेजी शब्दों से उनका स्टेम शब्द पृथक् करने के लिए 260 नियमों को प्रस्तावित किया था। पोर्टर (Porter) [2] ने विभिन्न अफिक्सो के आधार पर स्टेमर को विकसित किया था। इस स्टेमर की प्रक्रिया पाँच चरण में पूर्ण होती थी। पहले चरण में, इन्फ्लेक्शनल अफिक्स को नियंत्रित करते थे। अगले तीन चरण में, डेरिवेशनल सफिक्स को नियंत्रित किया एवं अंततः आखिरी चरण में रिकॉडिंग होती थी। यहाँ पर भी इन्फ्लेक्शनल मार्फोलोजी से तात्पर्य है कि स्टेम शब्द की श्रेणी उसके वास्तविक शब्द की श्रेणी के समान होती है एवं डेरिवेशनल मार्फोलोजी में स्टेम शब्द की श्रेणी उसके वास्तविक शब्द की श्रेणी से भिन्न होती है। खोजा एवं गर्सिदे (Khoja and Garside) [3] ने अरैबिक स्टेमर बनाया था जिसे सुपिरिओर रूट आधारित स्टेमर कहते थे। इस एल्गोरिथम में प्रेफिक्स, सफिक्स एवं इन्फिक्स तीनों उनके अफिक्स लिस्ट से मैप करा के पृथक् कर लेते थे एवं स्टेम शब्द को अलग प्रदर्शित कर देते थे। विशेषतः संज्ञा के साथ इसे विभिन्न समस्याओं का सामना करना पड़ता था। 20 वीं सदी के पश्चात,

शोधकर्ता मजूमदार एवं उनके साथियों (Majumder) [5] ने स्ट्रिंग्स के बीच की दूरी के आधार पर एक दृष्टिकोण पेश किया जो कि समूहों पर आधारित था और साथ ही इसे किसी भाषाई ज्ञान की आवश्यकता नहीं थी। इस दृष्टिकोण या प्रणाली को किसी भी भाषा के साथ लागू किया जा सकता है। अलशमारी एवं लिन (Al-Shammari and Lin) [5] ने एजुकेटेड टेक्स्ट स्टेमर का प्रस्तुतीकरण किया। यह बहुत ही सरल, शब्दकोश मुक्त एवं कुशल स्टेमर है जिसके द्वारा स्टेमिंग की त्रुटियों में गिरावट आई और साथ ही इसे कम स्टोरेज एवं कम प्रोसेसिंग टाइम की जरूरत पड़ती है। अकरम (Akram)[5] एवं उनके साथियों ने ऊर्दू भाषा के लिए एक स्टेमर प्रस्तुत किया। यह स्टेमर सिर्फ ऊर्दू भाषा के लिए प्रतिबंधित है यानि इस स्टेमर के द्वारा हम दूसरी किसी भाषा जैसे हिंदी, अंग्रेजी, फारसी के शब्दों से उनके स्टेम को पृथक् नहीं कर सकते हैं। इस स्टेमर के द्वारा हम प्रेफिक्स एवं सफिक्स दोनों को स्टेम शब्द से पृथक् करा सकते हैं। खान एवं उनके साथियों (Khan et. al.)[7] ने नियम- आधारित स्टेमर बनाने में आने वाली कई चुनौतियों को पेश किया है। इन लोगों ने अरबी, फारसी और ऊर्दू भाषा के लिए नियमों पर आधारित बने हुए स्टेमर की भी चर्चा की है। इस नियमों पर आधारित स्टेमर के द्वारा प्रेफिक्स एवं सफिक्स अलग किये गए हैं और इन्हें नियंत्रित करने के लिए कई पैटर्न निश्चित किये गये हैं। प्रेफिक्स वो होते हैं जो शब्द की शुरुआत में जुड़े होते हैं और सफिक्स शब्द के अंत में जुड़ा भाग होता है। मिश्रा एवं प्रकाश (Mishra and Prakash)[8] ने हिंदी भाषा के लिए एक प्रभावी स्टेमर का प्रस्ताव रखा। यह स्टेमर ओवर स्टेमिंग एवं अंडर स्टेमिंग की समस्या का निवारण करता है और इसके द्वारा 91.59 % की सटीकता प्राप्त होती है। हुसैन (Husain) [9] ने ऊर्दू और मराठी भाषा के स्टेमर के विकास के लिए एक अनसुपरवाईजड द्रष्टिकोण का प्रस्ताव रखा। इस द्रष्टिकोण में, उन्होंने सफिक्स प्रथकीकरण के दो तरीकों का प्रस्ताव किया : 1) फ्रीक्वेंसी /आवृत्ति

आधारित सफिक्स स्ट्रिपिंग एल्गोरिथम (Frequency based Suffix Stripping Algorithm) और 2) लम्बाई आधारित सफिक्स स्ट्रिपिंग एल्गोरिथम (Length based Suffix Stripping Algorithm)। जूही एवम् उनके साथियों (Juhi)[10] गुजराती – हिंदी मशीन अनुवाद की गुणवत्ता को बढ़ाने के लिये पोस (Part-of-Speech) टैगिंग एवं स्टेमिंग का इस्तेमाल किया है। इन्होंने 5400 पोस टैग आधारित वाक्यों पर 202 अफिक्स नियमों को लागू कर के अनुवाद किया तो इस प्रणाली की दक्षता 93.09 % मापी गई। स्निग्धा एवम् उनके साथियों (Snigdha)[11] ने नियमबद्ध हिंदी लेमेटाइजर को विकसित किया है। लेमेटाइजर के द्वारा हम स्टेमिंग में आने वाली कमियों को भी दूर कर सकते हैं। इसमें सफिक्स हटाने के साथ साथ कुछ नियम जोड़ने के लिये भी बनाये गये हैं। इस हिन्दी लेमेटाइजर की शुद्धता 89.08% मापी गई। मुबाशिर एवम् उनके साथियों (Mubashir) [12] ने रूपात्मक ऊर्दू की चुनौतियों को दूर करने के लिये एक प्रभावी प्रस्ताव दिया जिसमें नियम आधारित ऊर्दू स्टेमर की संरचना को दर्शाया गया है। यह स्टेमर ऊर्दू के स्टेम शब्द को उत्पन्न करता है और साथ ही कुछ दूसरी भाषाओं के स्टेम शब्दों को भी उत्पन्न करता है। जब्बर एवम् उनके साथियों (Jabbar) [13] ने ऊर्दू भाषा के लिये कई संसाधनों का विश्लेषण और विकास किया। साथ ही विभिन्न प्रकार की स्टेमिंग के तरीकों को भी परिभाषित किया हुआ है। 2019 में फिर से मुबाशिर एवम् उनके साथियों (Mubashir) [14] ने एक अन्य व्यापक स्टेमर का विकास किया जो कि खास तौर पर रूपात्मक तरीके से समृद्ध ऊर्दू भाषा के लिये बनाया गया है। यह स्टेमर नियमबद्ध अफिक्स हटाने की पद्धति पर आधारित है जिसके द्वारा स्टेम शब्द को प्राप्त किया जा सकता है।

3. ऊर्दू की भाषाई पृष्ठभूमि/ऊर्दू की मार्फोलोजी

ऊर्दू भाषा अन्य इन्डो-आर्यन भाषाओं के समान ही है। इसे दायीं से बायीं तरफ फारसी-अरबी लिपि के समान लिखा जाता है। ऊर्दू एक मुक्त शब्द क्रम

(free word order) एवं कमजोर विभक्ति वाली भाषा है। इसकी मार्फोलोजी बहुत ही जटिल एवं समृद्ध है। मार्फोलोजी का मतलब शब्दों की पूर्णतः बनावट अथवा आकृति विज्ञान से है। इसके द्वारा शब्दों को संज्ञा, क्रिया विशेषण इत्यादि शाब्दिक वर्गों में भी विभाजित किया जा सकता है। इसकी सबसे छोटी यूनिट 'मोर्फ़ीम' अथवा 'लीमा' कहलाती है। मोर्फ़ीम किसी भाषा की अर्थपूर्ण इकाई है जिसे आगे विभाजित नहीं किया जा सकता है। (Morpheme = free Phoneme = Lofue) उदाहरण के लिए शब्द – 'آدم' (आदम) एक एकल मोर्फ़ीम है किन्तु 'آدم زاد' (आदमजाद) शब्द में 'آدم' (आदम) एवं 'زاد' (जाद) दो मोर्फ़ीम हैं। इस उदाहरण से हम समझ सकते हैं कि कुछ शब्द उसके दो वर्ग : स्टेम तथा अफिक्स से मिलकर बने होते हैं। यहाँ 'स्टेम' ही किसी शब्द का मुख्य मोर्फ़ीम होता है। चूँकि अफिक्स के बजाय, स्टेम से ही हमें शब्द के अर्थ का ज्ञान हो पता है। यहाँ शब्दों को मोर्फ़ीम कि सहायता से बनाने के दो व्यापक तरीके हो सकते हैं : इन्प्लेक्शनल मोर्फ़ीम एवं डेरिवेशनल मोर्फ़ीम। इन्प्लेक्शनल मोर्फ़ीम में, वास्तविक शब्द एवं उसके स्टेम शब्द की श्रेणी समान होती है जैसे कि 'سوال + ات = سوالات' (सवाल+आत=सवालात) जहाँ 'سوال' (सवाल) एकवचन संज्ञा एवं 'سوالات' (सवालात) बहुवचन संज्ञा है। अर्थात् इनकी श्रेणी समान ही है। डेरिवेशनल मार्फ़ोलोजी में, वास्तविक शब्द एवं प्राप्त स्टेम शब्द या किसी शब्द में अफिक्स जोड़ने से प्राप्त नए शब्द की श्रेणी एक दूसरे से भिन्न होती है। उदाहरण के तौर पर – 'حکمت + ی = حکمتی' (हिकमत+ई=हिकमती), यहाँ शब्द की श्रेणी संज्ञा (اسم) से विशेषण (صفت) में परिवर्तित हो गई है।

ऊर्दू मोर्फ़ोलोजी में, संज्ञा, क्रिया, विशेषण एवं क्रिया-विशेषण के अनेक इन्प्लेक्शन रूप हो सकते हैं। इसके अंतर्गत लिंग, एकवचन, बहुवचन, काल इत्यादि भी शामिल हैं। हिंदी के समान ऊर्दू भी अनेक परसर्ग क्रिया एवं विशेषण के साथ जोड़ती है, जैसे कि – خوشی (खुशी), خوشحالی (खुशहाली), خشمجاز (खुशमिजाज), خوشنما (खुशनुमा), ناخوش (नाखुश)। ये

विभिन्न रूपांतरित शब्द “بخوش (खुश)” स्टेम के द्वारा बनाये गए हैं। अग्रित खंड में हम कुछ शाब्दिक वर्गों/श्रेणी जैसे اسم (संज्ञा), فیل (क्रिया), صفت (विशेषण) इत्यादि का संक्षिप्त विवरण प्रदर्शित कर रहे हैं।

A. संज्ञा (اسم): ऊर्दू में संज्ञा को तीन तरह से विभाजित कर सकते हैं।

1. केस : नोमिनेटीव (Nominative)/ओब्लिक (Oblique) /वोकेटीव (Vocative) (भाववाचक/अव्यक्त/शब्द वाचक)

2. नंबर : एकवचन /बहुवचन

3. लिंग : पुल्लिंग /स्त्रीलिंग

उदाहरण के लिए :-

- एकवचन पुल्लिंग संज्ञा सामान्यतः ۱ (अलिफ), ۲ (हे) एवं ۳ (ऐन) के द्वारा समाप्त होती है।
- एकवचन स्त्रीलिंग बनाने के लिए अंत में ۴ (इ) जोड़ देते हैं।
- बहुवचन नोमिनेटीव एवं एकवचन ओब्लिक बनाने के लिए आखिरी अक्षर को ۵ (ये) से प्रतिस्थापित किया जाता है।
- बहुवचन ओब्लिक बनाने के लिए, आखिरी अक्षर को ۶ (ओं) मोर्फीम से प्रतिस्थापित किया जाता है।
- बहुवचन वोकेटीव बनाने के लिए आखिरी अक्षर को ۷ (वोव) से प्रतिस्थापित करते हैं।

तालिका 1- : संज्ञा समूह का उदाहरण

	नोमिनेटिव	ओब्लिक	वोकेटिव
एकवचन	بہودا (हथोड़ा)	بہودے (हथौड़े)	بہودے (हथौड़े)
बहुवचन	بہودے (हथौड़े)	بہودوں (हथौड़ों)	بہاوڈو (हथौड़ो)

B. क्रिया (فیل): ऊर्दू क्रिया दूसरे शाब्दिक वर्गों कि तुलना में बहुत जटिल है। मार्फोलोजी के

सन्दर्भ में, ऊर्दू क्रिया के अनेक रूप हैं –

- लिंग : पुल्लिंग, स्त्रीलिंग
- नंबर : एकवचन, बहुवचन
- व्यक्ति : प्रथम, द्वितीय एवं तृतीय

ऊर्दू क्रिया प्रत्यक्ष एवं अप्रत्यक्ष प्रेरणा का व्यवहार प्रस्तुत करती है। आम तौर पर क्रिया किसी “स्टेम” से ही बनती है। उदाहरण के लिए यदि हम क्रिया ‘कर’ की बात करें तो—

- साधारण रूप : करना
- प्रत्यक्ष साधारण रूप : कराना
- अप्रत्यक्ष साधारण रूप : करवाना

ये तीनों क्रियाएँ, क्रिया ‘कर’ की शाब्दिक विच्छेद हैं।

C. विशेषण (صفت): ऊर्दू विशेषण भी लिंग, नंबर, एवं कारक के आधार पर इन्फ्लेक्टेड (विभक्त) होते हैं। उदाहरण के लिए – فریلا (फुर्तीला) एकवचन पुल्लिंग रूप है, जिसे यदि बहुवचन पुल्लिंग में परिवर्तन करना हो तो فریله (फुर्तीले) में परिवर्तित किया जाएगा। यदि अब हम स्त्रीलिंग में परिवर्तित करना चाहे तो हमें فریلی (फुर्तीली) लिखना होगा।

D. क्रिया-विशेषण (فیل-صفت) : क्रिया-विशेषण, विशेषण की तरह ही परिवर्तिनीय एवं अपरिवर्तिनीय रूप में होती है। यह संज्ञा एवं क्रिया के रूप परिवर्तन पर आधारित है। हम क्रिया-विशेषण को कुछ निम्नलिखित वर्गों में विभाजित कर सकते हैं।

- समय की क्रिया-विशेषण : روجانا (रोजाना) / اکثر (अक्सर)
- स्थान की क्रिया-विशेषण : وہاں (यहाँ) / یہاں (वहाँ)
- व्यवहार की क्रिया-विशेषण : یکایک (यकायक)
- कोटि की क्रिया-विशेषण : چھوٹا (छोटा) / لمبا (लम्बा)

4. प्रस्तावित प्रणाली

इस प्रस्तावित प्रणाली में हमने नियमबद्ध ऊर्दू स्टेमर का निर्माण किया है। इस नियमबद्ध स्टेमर के द्वारा, इन्प्लेक्शनल एवं डेरिवेशनल दो तरह के मोर्फ़ीम, "स्टेम" के तौर पर प्राप्त होते हैं। कुछ इन्प्लेक्शनल एवं डेरिवेशनल मोर्फ़ोलोजी को प्रदर्शित करने वाले शब्द नीचे की तालिका 2 एवम् 3 में दर्शाये गए हैं जो कि हमारे नियमबद्ध स्टेमर के द्वारा प्राप्त हुए हैं।

तालिका 2 : इन्प्लेक्शनल मोर्फ़ीम /स्टेम

शब्द	अफिक्स	स्टेम
برسات (बरसात) — संज्ञा	ات (आत)	برس (बरस) — संज्ञा
حمایتی (हिमायती) — संज्ञा	ی (ई)	حمایت (हिमायत) — संज्ञा
حلقنامه (हलफनामः) — संज्ञा	نام (नामः)	حلف (हलफ) — संज्ञा

तालिका 3 : डेरिवेशनल मोर्फ़ीम /स्टेम

शब्द	अफिक्स	स्टेम
آبادی (आबादी) — संज्ञा	ی (ई)	آباد (आबाद) — विशेषण
بے حیا (बेहया) — विशेषण	بے حیا (बे)	حیا (हया) — संज्ञा
خاردار (खारदार) — विशेषण	دار (दार)	خار (खार) — संज्ञा

प्रस्तावित प्रणाली को विकसित करने के लिये हम निम्नलिखित चरणों का पालन करेंगे, जो कि इस तरह से हैं —

4.1 डेटा संग्रह :

ऊर्दू डेटा संग्रह के लिये हमने पर्यटन एवम् स्वास्थ्य विभाग के डेटा को संग्रहित किया है, जिसे निम्नलिखित तालिका 4 में प्रदर्शित किया गया है —

तालिका 4: विस्तारित डेटा संग्रह

डेटासेट	वाक्यों की संख्या	शब्दों की संख्या	3 या 3 से बड़ी लंबाई वाले शब्द
पर्यटन	25000	441617	338745
स्वास्थ्य	25000	441197	339571
कुल	50000	882814	678316

4.2 अफिक्स उत्पन्न :

स्टेमर के विकास के लिए हमने बहुत सारे शब्दों का निरीक्षण किया। इस निरीक्षण से हमें ज्ञात हुआ कि किसी भी शब्द से कुछ प्रेफिक्स एवं सफिक्स जोड़ने एवं हटाने से उनके शब्द विच्छेद में किस किस तरह के परिवर्तन आते हैं। यहाँ हमने 678316 शब्दों के निरीक्षण के पश्चात बिना दोहराये हुए 146 अफिक्स की सूची नियमों के तौर पर तैयार की है, जिनमें से 129 सफिक्स एवं 17 प्रेफिक्स प्राप्त हुए हैं।

ये अफिक्स किसी भी शब्द से स्टेम को पृथक करने के लिए इस्तेमाल किये गए हैं। नीचे के चित्र 1 में कुछ अफिक्सों को दर्शाया गया है।

انیں	یں	ی	و	ے
نے	نے	ین	ات	ا
انی	وں	یاں	تا	یوں
ں	بد	نو	ناک	گے

चित्र 1 : नमूना आधारित अफिक्स सूची

उपरोक्त सूची में कई अफिक्स नहीं भी शामिल किया है क्योंकि कुछ ऐसे भी अफिक्स हैं जिनको

शामिल करने की वजह से ओवरस्टेमिंग की समस्या का सामना करना पड़ सकता है। इस नियमबद्ध प्रणाली में हम दो तरीके से अफिक्स को स्टेम शब्द से हटा सकते हैं।

4.2.1 प्रणाली 1 : सबसे बड़े प्रत्यय के आधार पर प्रत्यय निष्कासन (Largest suffix based affix removal algorithm)

इस प्रणाली में हम अफिक्स को उसकी लम्बाई के आधार पर घटते क्रम में रखते हैं। यह प्रणाली, स्ट्रिपिंग पद्धति पर आधारित है जो कि अफिक्स को उसके स्टेम शब्द से हटाती जाती है। बड़ी लम्बाई वाले अफिक्स पहले हटेंगे फिर यदि जरूरत होगी तो उससे छोटी लम्बाई वाले अफिक्स हटेंगे। कुछ शब्दों में से हमें उचित स्टेम नहीं भी प्राप्त होता है। उदाहरण के लिए : 'حيات' (हयात) शब्द से कुछ भी नहीं हटना चाहिए किन्तु जब इस शब्द को अपने स्टेम से प्रोसेज कराएँगे तो हमें 'بے' (ह्य) स्टेम एवं सफिक्स 'ات' (आत) अलग हो जाएगा। यह स्टेम कोई अर्थ नहीं रखता है। ऐसे अपवाद शब्दों के लिए हमने एक डेटाबेस तैयार किया है एवम् सभी अपवाद शब्दों को इसी डेटाबेस में ही समाहित किया हुआ है।

4.2.2 प्रणाली 2 : उच्चतम आवृत्ति के आधार पर प्रत्यय निष्कासन (Highest frequency based affix removal algorithm)

इस प्रणाली में हमने अपने शब्दों से उत्पन्न अफिक्सों की उनके घटते हुए आवृत्ति के आधार पर एक सूची बनाई। इस सूची में उच्चतम आवृत्ति वाले अफिक्सों को सबसे ऊपर और कम आवृत्ति वालों को सूची में सबसे नीचे रखा है, जिसके अंतर्गत उच्चतम आवृत्ति वाले अफिक्स सबसे पहले शब्दों से हटेंगे और फिर उसका स्टेम उत्पन्न करेंगे।

4.3 डेटाबेस :

यह डेटाबेस हमने ऊर्दू के कुछ ऐसे शब्दों के लिए बनाया है जिनमे से कोई भी अफिक्स हटाने की आवश्यकता नहीं होती है। अर्थात यदि हम इन

शब्दों से कोई भी अफिक्स हटाते हैं तो उनसे मिले स्टेम का कोई अर्थ नहीं रह जाता है। इस डेटाबेस की सहायता से ओवर-स्टेमिंग की समस्या भी कम हुई है। डेटाबेस में रखे हुए अपवाद शब्दों में से कुछ नीचे के चित्र 2 में दर्शाये गए हैं।

لڑکا (लड़का)	لڑکی (लड़की)	حوالات (हवालात)
خرافات (खुराफात)	لڑکے (लड़के)	مرتب (मरतबा)
رستمالا (रस्तमाला)	سپنا (सपना)	شادی (शादी)
سپرا (सपेरा)	نخرا (नखरा)	دھوکھے (धोखे)

चित्र 2 : अपवाद शब्द

4.4 एल्गोरिथम :

नियमबद्ध ऊर्दू स्टेमर के लिए निम्नलिखित एल्गोरिथम को प्रस्तावित किया गया है।

स्टेप-1 : इनपुट ग्रहण कराएँ।

स्टेप-2 : इनपुट चेक करें। इनपुट कोई वाक्य या एकल शब्द हो सकता है।

स्टेप-3 : यदि इनपुट कोई वाक्य है तो उसे शब्दों में टोकेनाइज (विभाजित) करें।

स्टेप-4 : हर एक शब्द (शब्दों की लंबाई 3 या 3 से बड़ी होनी चाहिये) के लिए निम्नलिखित स्टेप का अनुसरण करें।

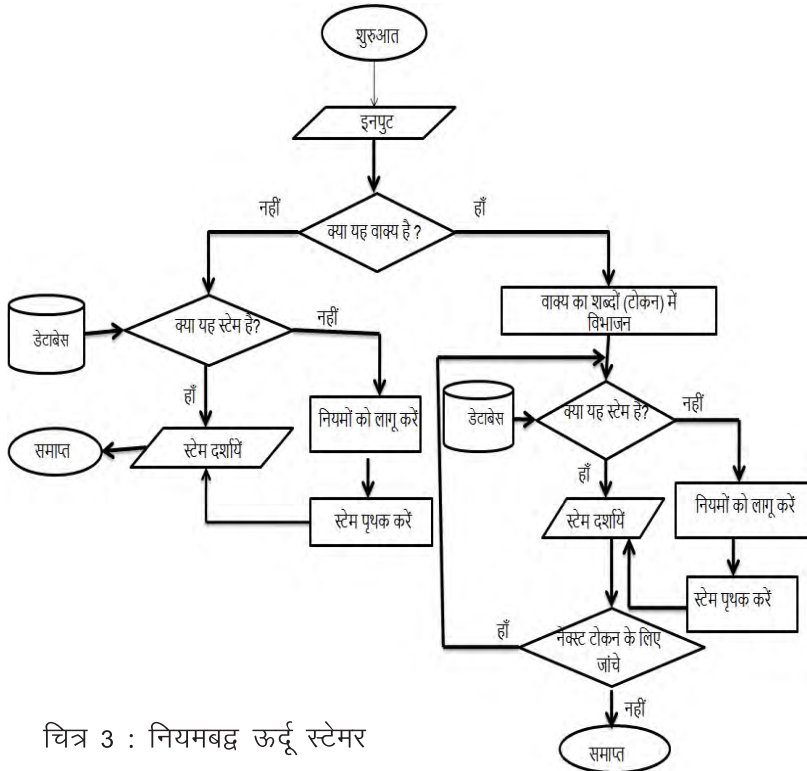
- डेटाबेस से मैच कराये यदि इनपुट डेटाबेस में उपलब्ध है तो स्टेम के समान उसी इनपुट को प्रदर्शित करा देंगे।
- यदि इनपुट डेटाबेस में उपलब्ध नहीं है तो अफिक्स सूची के नियमों को प्रणाली 1 अथवा प्रणाली 2 के आधार पर लागू करा के स्टेम शब्द प्राप्त करेंगे और फिर इस स्टेम को प्रदर्शित कराएँगे।

स्टेप-5 : यदि इनपुट कोई शब्द है तो उपर्युक्त स्टेप- 4 का अनुसरण करें।

अब नीचे की तालिका-5 में इस स्टेमिंग एल्गोरिथम से प्राप्त कुछ परिणाम को प्रदर्शित किया है, आगे हम इस स्टेमिंग एल्गोरिथम को फ्लोचार्ट (चित्र 3) के द्वारा प्रदर्शित कर रहे हैं।

शब्द	स्टेम	प्रेफिक्स	सफिक्स
حضورى (हुजूरी)	حضور (हुज़ूर)		ی (ई)
سوالات (सवालात)	سوال (सवाल)		ات (आत)
نظرے (नज़रे)	نظر (नज़र)		ے (ए)
مجلسی (मजलसी)	مجلس (मजलिस)		ی (ई)
بھولا پن (भोलापन)	بھولا (भोला)		پن (पन)
ایماندار (ईमानदार)	امان (ईमान)		دار (दार)
نوجوان (नौजवान)	جوان (जवान)	نو (नौ)	
لا جواب (लाजवाब)	جواب (जवाब)	لا (ला)	
بدنصیب (बदनसीब)	نصیب (नसीब)	بد (बद)	
بے ادب (बेअदब)	ادب (अदब)	بے (बे)	

तालिका - 5:
ऊर्दू स्टेमर द्वारा प्राप्त परिणाम



चित्र 3 : नियमबद्ध ऊर्दू स्टेमर

5. प्रणाली का परिणाम एवम् मूल्यांकन

इस प्रणाली को बनाने के साथ साथ अब हम यह जानना चाहते हैं कि यह प्रणाली कितने प्रतिशत शुद्ध परिणाम देती है अर्थात् इस प्रणाली का शुद्धता के आधार पर मूल्यांकन करना चाहते हैं। शुद्धता ज्ञात करने के लिए निम्नलिखित सूत्र का इस्तेमाल किया गया है।

$$\text{शुद्धता } (:) = (\text{शुद्ध प्राप्त स्टेम शब्द}) / (\text{कुल इनपुट शब्द}) \times 100$$

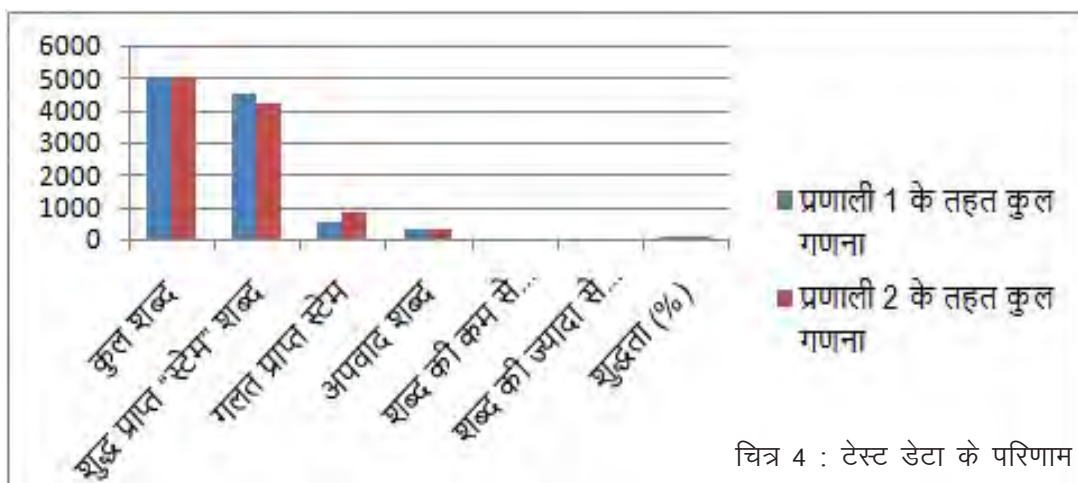
इस शोध पत्र में हमने अपनी प्रणाली की जांच के लिए 5000 शब्दों पर यह एल्गोरिथम लगाई। इन शब्दों को हमने ऊर्दू के लेख से एकाएक चयन किया है। इस जांच वाले डेटा का संक्षिप्त विवरण नीचे की तालिका में दर्शाया गया है।

प्रणाली 1 एवम् प्रणाली 2 के अंतर्गत डेटा का मूल्यांकन करने के लिये हमने इन्हें अपनी एल्गोरिथम में लागू किया। इसके पश्चात हमें कुछ इस प्रकार से परिणाम प्राप्त होते हैं जो कि तालिका 6 में दर्शाये गये हैं –

तालिका 6 : प्रणाली 1 एवम् प्रणाली 2 पर आधारित जांच डेटा का संक्षिप्त विवरण

जांच डेटा के लक्षण	प्रणाली 1 के अंतर्गत कुल गणना	प्रणाली 2 के अंतर्गत कुल गणना
कुल शब्द	5000	5000
शुद्ध प्राप्त "स्टेम" शब्द	4483	4203
गलत प्राप्त स्टेम	517	797
अपवाद शब्द	312	312
शब्द की कम से कम लम्बाई	3	3
शब्द की ज्यादा से ज्यादा लम्बाई	15	15
शुद्धता (%)	89.6	84.06

उपर्युक्त सूत्र के अनुसार, यदि इस जांच डेटा के द्वारा अपनी प्रणालियों की मैन्युअली (manually) शुद्धता मापते हैं तो हमें 89.60 % एवम् 84.06 % की शुद्धता प्राप्त होती है। इन परिणामों से हम यह अनुमान लगा सकते हैं कि प्रणाली 1 जो कि सबसे बड़े प्रत्यय के आधार पर प्रत्यय निष्कासन करती है वो ज्यादा अच्छा एवम् शुद्ध परिणाम प्रदान करती है। नीचे दिया गया चार्ट (चित्र 4) भी उपर्युक्त टेस्ट डेटा के परिणामों को प्रदर्शित करता है।



चित्र 4 : टेस्ट डेटा के परिणाम

6. निष्कर्ष

इस शोध पत्र में हमने नियमबद्ध ऊर्दू स्टेमर के विकास को दर्शाया है और साथ ही कुछ अपवाद शब्दों के लिए डेटाबेस तैयार किया है, जिससे ओवरस्टेमिंग की समस्या कम हुई है। यह नियमबद्ध स्टेमर सिर्फ ऊर्दू भाषा के लिए ही बनाया गया है। इसके द्वारा हम किसी भी शब्द के स्टेम को पृथक कर सकते हैं। इस स्टेमर की शुद्धता प्रणाली 1 के अंतर्गत 89.60 % और प्रणाली 2 के अंतर्गत 84.06 % मापी गई है। भविष्य में हम इस पर और कार्य करके इसकी अफिक्स सूची को और बढ़ाएंगे। इस प्रक्रिया को हम किसी दूसरी भाषा के लिये भी इस्तेमाल कर सकते हैं। किसी अन्य भाषा के लिये हमें अनिवार्य रूप से उनकी अफिक्स सूची और उसी भाषा के अपवादों की सूची तैयार करनी होगी।

Table of English Terminology which is used in paper:

English	Hindi
Algorithm	कलन विधि
Analysis	विश्लेषण
Derivational Morphology	व्युत्पन्न आकृति विज्ञान
Flowchart	प्रवाह संचित्र
Inflectional Morphology	अनिमेश आकृति विज्ञान
Longest Suffix Stripping	सबसे लंबा प्रत्यय अलग करना
Prefix	उपसर्ग
Process	प्रक्रिया
Removal	निष्कासन
Spell Checker	वर्तनी परीक्षक
Suffix	प्रत्यय

सन्दर्भ

- [1] Julie Beth Lovins, Development of Stemming Algorithm, Mechanical Translation and Computational Linguistics, Vol. 11, No. 1, pp 22-31, 1968
- [2] M.F. Porter, An algorithm for suffix stripping, Program, 14 (3) pp. 130-137. 1980.
- [3] S. Khoja and R. Garside. Stemming Arabic Text. Lancaster, UK, Computing Department, Lancaster University. 1999.
- [4] P. Majumder, M. Mitra, S. K. Parui, G. Kole, P. Mitra and K. Datta. "YASS: Yet another suffix stripper", ACM Transactions on Information Systems, Vol. 25, No. 4, pp. 18-38. 2007.
- [5] Eiman Tamah Al-Shammari, Jessica Lin, Towards an Error-Free Arabic Stemming, iNEWS'08, Napa Valley, California, USA. 2008.
- [6] Q. Akram, A. Naseer and S. Hussain, As-sas-Band, an Affix- Exception-List Based Urdu Stemmer, In Proceedings of the 7th Workshop on Asian Language Resources, pp. 40- 47, Suntec, Singapore. 2009.
- [7] Sajjad Ahmad Khan, Waqas Anwar, Usama Ijaz Bajwa, Challenges in Developing a Rule based Urdu Stemmer. In Proceedings of the 2nd Workshop on South and Southeast Asian Natural Language Processing (WSSANLP), IJCNLP 2011., pages 46-51, Chiang Mai, Thailand. 2011.
- [8] Upendra Mishra and Chandra Prakash, MAULIK: An Effective Stemmer for Hindi Language. International Journal on Computer Science and Engineering (IJCSE), pp. 711-717. 2012.
- [9] Mohd. Shahid Husain. An Unsupervised Approach To Develop Stemmer. International Journal on Natural Language Computing (IJNLC) Vol. 1, No.2 2012.

- [10] J.Ameta, N.Joshi, I.Mathur. Improving the quality of Gujarati Hindi Machine Translation Through Part-of-Speech Tagging and Stemmer-Assisted Transliteration. In proceeding of *International Journal on Natural Language Computing, Vol 3 (2), pp 49-54, 2013.*
- [11] S.Paul, M.Tandon, N.Joshi, I.Mathur. Design of a Rule Based Hindi Lemmatizer. In proceeding of the *Third International workshop on Artificial Intelligence, Soft Computing And Application, Chennai, India, pp 67-74, 2013.*
- [12] Ali M., Khalid S., and Saleemi M., "A Novel Stemming Approach for Urdu language," *Journal of Applied Environmental and Biological Sciences*, vol. 4, no. 7S, pp. 436-443, 2014.
- [13] Jabbar Abdul, Sajid Iqbal, and Muhammad Usman Ghani Khan. "Analysis and development of resources for Urdu text stemming." *Language and Technology* 1. 2016.
- [14] Mubashir Ali, Shehzad Khalid, and Muhammad Saleemi. Comprehensive Stemmer for Morphologically Rich Urdu Language. *The International Arab Journal of Information Technology*, Vol. 16, No. 1, pp-139-147, January 2019.

प्रेरक प्रसंग - दीप से दीप जलाओ

ईश्वर चन्द्र विद्यासागर (1820-1891) उन्नीसवीं शताब्दी के बंगाल के प्रसिद्ध दार्शनिक, शिक्षाविद, समाज सुधारक, लेखक, अनुवादक, मुद्रक, प्रकाशक, उद्यमी और परोपकारी व्यक्ति थे। वे बंगाल के पुनर्जागरण के स्तम्भों में से एक थे। उनके बचपन का नाम **ईश्वर चन्द्र बन्दोपाध्याय** था। संस्कृत भाषा और दर्शन में अगाध पाण्डित्य के कारण विद्यार्थी जीवन में ही संस्कृत कॉलेज ने उन्हें 'विद्यासागर' की उपाधि प्रदान की थी। ईश्वर चन्द्र विद्यासागर कलकत्ता में अध्यापन कार्य करते थे। वेतन का उतना ही अंश घर परिवार के लिए खर्च करते जितने में कि औसत नागरिक स्तर का गुजारा चल जाता। शेष भाग वे दूसरे जरूरतमंदों की, विशेषतया छात्रों की सहायता में खर्च कर देते थे। आजीवन उनका यही व्रत रहा।

एक दिन वे बाजार में चले जा रहे थे। एक हताश युवक ने भिखारी की तरह उनसे एक पैसा माँगा। विद्यासागर दानी तो थे पर सत्पात्र की परीक्षा किये बिना किसी की ठगी में न आते। युवक से जवानी में हट्टे कट्टे होते हुए भी भीख माँगने का कारण पूछा। सारी स्थिति जानने पर माँगने का औचित्य लगा। सो एक पैसा तो दे दिया पर उसे रोककर उससे पूछा कि यदि अधिक मिल जाय तो क्या करोगे? युवक ने कहा कि यदि एक रुपया मिले तो उसका सौदा लेकर गलियों में फेरी लगाने लगूंगा और अपने परिवार का पोषण करने में स्वावलम्बी हो जाऊंगा।

विद्यासागर ने एक रुपया उसे और दे दिया। उसे लेकर उसने छोटा व्यापार आरंभ कर दिया। काम दिनों-दिन बढ़ने लगा। कुछ दिन में वह बड़ा व्यापारी बन गया।

एक दिन विद्यासागर उस रास्ते से निकल रहे थे कि व्यापारी दुकान से उतरा, उनके चरण स्पर्श किए, और उन्हें दुकान दिखाने ले गया, और कहा - यह आपकी दी गयी एक रुपये की पूंजी का चमत्कार है। विद्यासागर प्रसन्न हुए और कहा, जिस प्रकार तुमने सहायता प्राप्त करके उन्नति की उसी प्रकार का लाभ अन्य जरूरतमंदों को भी देते रहना। व्यापारी ने वैसा ही करते रहने का वचन दिया।

इस दृष्टांत का संदेश यह है कि उत्साही युवा उद्यमियों को नवाचारमय उद्योग सृजन में प्रोत्साहन पूँजी की व्यवस्था से आत्मनिर्भरता की क्रान्ति संभव है।

वायुमंडलीय कणिकीय द्रव्यों PM_{2.5} एवं PM₁₀ के स्तरों पर लॉकडाउन का प्रभाव: नई दिल्ली को लेकर एक केस अध्ययन

Effect of lockdown on the levels of atmospheric particulate matter PM_{2.5} and PM₁₀: A case study from New Delhi

डॉ० शिल्पी शाक्य¹ एवं डॉ० बिन्ध्याचल यादव²

¹सहायक प्राध्यापक, जन्तु विज्ञान विभाग, राजकीय महाविद्यालय फतेहाबाद, आगरा

²सहायक प्राध्यापक, वनस्पति विज्ञान विभाग, राजकीय महाविद्यालय फतेहाबाद, आगरा

¹shilpy.shilpy@gmail.com, ²bindhyachal@gmail.com

सारांश

कोरोनावायरस (COVID-19) सर्वव्यापी महामारी ने सम्पूर्ण मानव जाति को प्रभावित कर दिया है। पूरी दुनिया के साथ भारत भी पिछले कुछ महीनों से इस सर्वव्यापी महामारी का दंश झेल रहा है। कोरोनावायरस संक्रमण के प्रसार को नियंत्रित करने के उद्देश्य से देश भर में 25 मार्च 2020 को पूर्ण लॉकडाउन घोषित किया गया। लॉकडाउन के कारण उत्पन्न हुई कई मुश्किलों के बीच अप्रत्यक्ष रूप से हमारे वातावरण ने सकारात्मक पुनर्प्राप्ति दर्शायी। शीर्ष भारतीय शहर जिनका दुनियाभर में सबसे खराब हवा की गुणवत्ता का ट्रैक रिकॉर्ड रहा है, लॉकडाउन के दौरान उनकी हवा की गुणवत्ता में भी एक अप्रत्याशित सुधार देखा गया। सरकार द्वारा लगाये गए लॉकडाउन का द्वितीयक प्रभाव देश के कई क्षेत्रों में पड़ा जहाँ वायु प्रदूषण के स्तर में महत्वपूर्ण कमी दर्ज की गयी। प्रस्तुत शोधपत्र में हमने पिछले एक वर्ष के दौरान नयी दिल्ली स्थित अशोक विहार के रिहाइशी इलाकों में चुने हुए प्रदूषकों (PM_{2.5} एवं PM₁₀) के स्तर का वैज्ञानिक विश्लेषण किया है। इसके साथ लॉकडाउन की अवधि व अनलॉक के उपरान्त, वायु में शामिल अन्य महत्वपूर्ण प्रदूषकों (PM_{2.5}, PM₁₀, NO₂, NH₃, CO, एवं CO₂) के स्तर का भी अध्ययन किया है। इस अध्ययन में मात्रात्मक विश्लेषण के आधार पर हम निम्नलिखित निष्कर्ष पर पहुंचे (1) मानव-संबंधी गतिविधियाँ वायु गुणवत्ता के साथ जुड़ी हुई हैं। (2) वायु प्रदूषण में कमी इस सर्वव्यापी महामारी के दौरान यात्रा प्रतिबंधों के साथ जुड़ी पायी गयी – औसतन, वायु गुणवत्ता सूचकांक (Air Quality Index) में लॉकडाउन से पूर्व और लॉकडाउन के दौरान 35.33 प्रतिशत की कमी पाई गयी और चयनित वायु प्रदूषक (PM_{2.5} एवं PM₁₀) में लॉकडाउन से पूर्व और लॉकडाउन के दौरान क्रमशः 40.7 प्रतिशत एवं 41.60 प्रतिशत की कमी दर्ज की गयी। (3) हमारे निष्कर्ष हरित उत्पादन और उपभोग की भूमिका को समझने के महत्व पर प्रकाश डालते हैं, जैसा कि स्पष्ट है, लॉकडाउन के परिणामस्वरूप वायु प्रदूषण में आई इस अस्थायी गिरावट को हम अनलॉक के दौरान बनाये रखने में सफल नहीं हुए हैं, परन्तु इस कदम से यह साबित हो गया है कि लॉकडाउन जैसे कठोर कदम भविष्य में आकस्मिक विकल्प के रूप में वायु प्रदूषण के नियंत्रण एवं उसके परिणाम स्वरूप होने वाले नुकसान से मानव जाति को बचा सकते हैं।

ABSTRACT

The coronavirus (COVID-19) pandemic has affected the entire human race. India as well as the entire world has been bearing the brunt of this pandemic for the last few months. In order to control

the spread of coronavirus infection, a complete lockdown was declared on March 25, 2020 across the country. Our environment showed a positive recovery indirectly amidst many of the difficulties caused by the lockdown. The top Indian cities, which have a track record of the worst air quality worldwide, also showed an unexpected gain in their air quality amid lockdown. The government-imposed lockdown had a secondary effect in many areas of the country where significant reduction in air pollution level was recorded. In the present paper, we have done a scientific analysis of the level of selected pollutants (PM2.5 and PM10) in the residential areas of Ashok Vihar in New Delhi during the last one year. In addition, the levels of other important pollutants (PM2.5, PM10, NO₂, NH₃, CO, and CO₂) in the air during the lockdown and post lockdown have also been studied. Based on the quantitative analysis of this study, we came to the following conclusions (1) Human-related activities are associated with the air quality. (2) Reduction in air pollution was associated with travel restrictions during this pandemic - on average, the Air Quality Index was found to be 35.33 percent lower before lockdown and during lockdown and selected air pollutants (PM2.5 and PM10) recorded a decrease of 40.7 percent and 41.60 percent before lockdown and during lockdown respectively. (3) Our findings emphasize the importance of the role of green production and consumption. It is clear that we have not been able to sustain this temporary decline in air pollution as a result of the lockdown during the periods of unlock, but this step has proved that drastic steps like lockdown as a contingency option in the future for controlling air pollution and saving mankind from consequent loss.

मुख्य शब्द : कोविड-19, लॉकडाउन, PM2.5, PM10, प्रदूषण, राष्ट्रीय वायु सूचकांक

Key words: COVID-19, Lockdown, PM2.5, PM10, Pollution, National air quality index

प्रस्तावना

प्रदूषण का तात्पर्य पृथ्वी के पर्यावरण के प्रदूषण से है जो मानव स्वास्थ्य, जीवन की गुणवत्ता एवं पारिस्थितिक तंत्र के कार्य में बाधा डालते हैं। प्रदूषण के प्रमुख रूपों में जल प्रदूषण, वायु प्रदूषण, ध्वनि प्रदूषण और मृदा (soil) प्रदूषण शामिल हैं। दिल्ली या राष्ट्रीय राजधानी क्षेत्र (National Capital Territory) केन्द्र और राज्य सरकारों द्वारा संयुक्त रूप से प्रशासित किया जाता है। इसमें लगभग 11.03 करोड़ लोग रहते हैं [1]। नई दिल्ली दुनिया का दूसरा सबसे बड़ा महानगर है। यह शहर भारत में शहरी आबादी का सबसे बड़ा योगदानकर्ता है। नई दिल्ली की आबादी औसत 37.60% की वार्षिक दर से बढ़ रही है। नई दिल्ली में अनियंत्रित शहरी विकास, बढ़ता हुआ औद्योगीकरण एवं अभूतपूर्व आबादी के प्रभाव ने पर्यावरण प्रदूषण को गंभीर रूप दिया है। दिल्ली के वायु प्रदूषण में वाहनों से होने वाले प्रदूषण का महत्वपूर्ण योगदान है। इसका सबसे स्पष्ट परिणाम वायु की गुणवत्ता का बिगड़ना है। वायु प्रदूषण मनुष्यों में होने वाले 40% इस्केमिक हृदय रोग, 40% स्ट्रोक, 11% पुरानी प्रतिरोधी फुफ्फुसीय रोग, 6% फेफड़े का कैंसर एवं 3% बच्चों में तीव्र श्वसन संक्रमण (Acute respiratory infection) के लिए एक महत्वपूर्ण एवं जिम्मेदार कारक है। दिल्ली के वायु प्रदूषण और मृत्यु दर पर किए गए अध्ययन में पाया गया कि बढ़े हुए वायु प्रदूषण के साथ मृत्यु दर और रुग्णता (Morbidity) बढ़ गई [2]। वायु प्रदूषण में PM2.5 एवं PM10 मानक का उपयोग आमतौर पर वायु की गुणवत्ता को मापने के लिए किया जाता है। PM10 मानक में 10 माइक्रोन या उससे कम व्यास (0.0004 इंच या एक मानव बाल की चौड़ाई का सातवां भाग) के कण शामिल हैं [3]। PM2.5 ऐसे वायुमंडलीय कणों को संदर्भित करता है जिनका

व्यास 2.5 माइक्रोमीटर से कम होता है। यह हमारे बालों के व्यास का लगभग 3 प्रतिशत है। PM2.5 श्रेणी के कण इतने छोटे होते हैं कि उन्हें केवल इलेक्ट्रॉन माइक्रोस्कोप की मदद से देखा जा सकता है। श्वसन पथ (Respiratory Tracts) के निचले क्षेत्रों तक पहुंचने की उनकी क्षमता के कारण इन छोटे कणों के स्वास्थ्य पर प्रतिकूल प्रभाव पड़ने की संभावना है [4, 5]। पूरी दुनिया के साथ-साथ भारत में भी पिछले तीन महीनों से कोरोनावायरस (COVID-19) सर्वव्यापी महामारी ने देश भर की आबादी के जीवन को प्रभावित किया। कोरोनावायरस संक्रमण के प्रसार को धीमा करने के प्रयास में देश भर के सभी वाणिज्यिक और औद्योगिक गतिविधियों पर प्रतिबंध लगाने के साथ पूर्ण लॉकडाउन किया गया। कोरोनावायरस के संक्रमण को फैलने से रोकने हेतु सम्पूर्ण भारत में 25 मार्च 2020 को देशव्यापी लॉकडाउन लगाया गया। वैश्विक स्तर पर इस तरह का कदम मानव इतिहास में सबसे पहला एवं सबसे बड़ी संगरोध विधि (Quarantine) के रूप में देखा गया। देशव्यापी लॉकडाउन लगाकर मानव गतिशीलता, उत्पादन और खपत गतिविधियों पर रोक लगाए जाने के पीछे के वैज्ञानिक कारणों में से सबसे महत्वपूर्ण था कोरोना विषाणु को संक्रमित व्यक्तियों से स्वस्थ व्यक्तियों में प्रसारित होने से रोकना। लॉकडाउन के कारण उत्पन्न कई मुश्किलों के बीच अप्रत्यक्ष रूप से हमारे वातावरण एवं प्राणवायु की गुणवत्ता में काफी सुधार हुआ है। सरकार के इस कदम का द्वितीयक प्रभाव (Secondary Effects) देश के कई क्षेत्रों में पड़ा जहाँ वायु प्रदूषण के स्तर में महत्वपूर्ण कमी दर्ज की गयी। प्रस्तुत शोध-पत्र में अशोक विहार, नई दिल्ली में वायु प्रदूषण की स्थिति को लॉकडाउन के पूर्व, लॉकडाउन के दौरान एवं पिछले एक वर्ष की स्थिति को एक साक्ष्य-आधारित अध्ययन के माध्यम से प्रस्तुत किया गया है।

सामग्री और तरीके

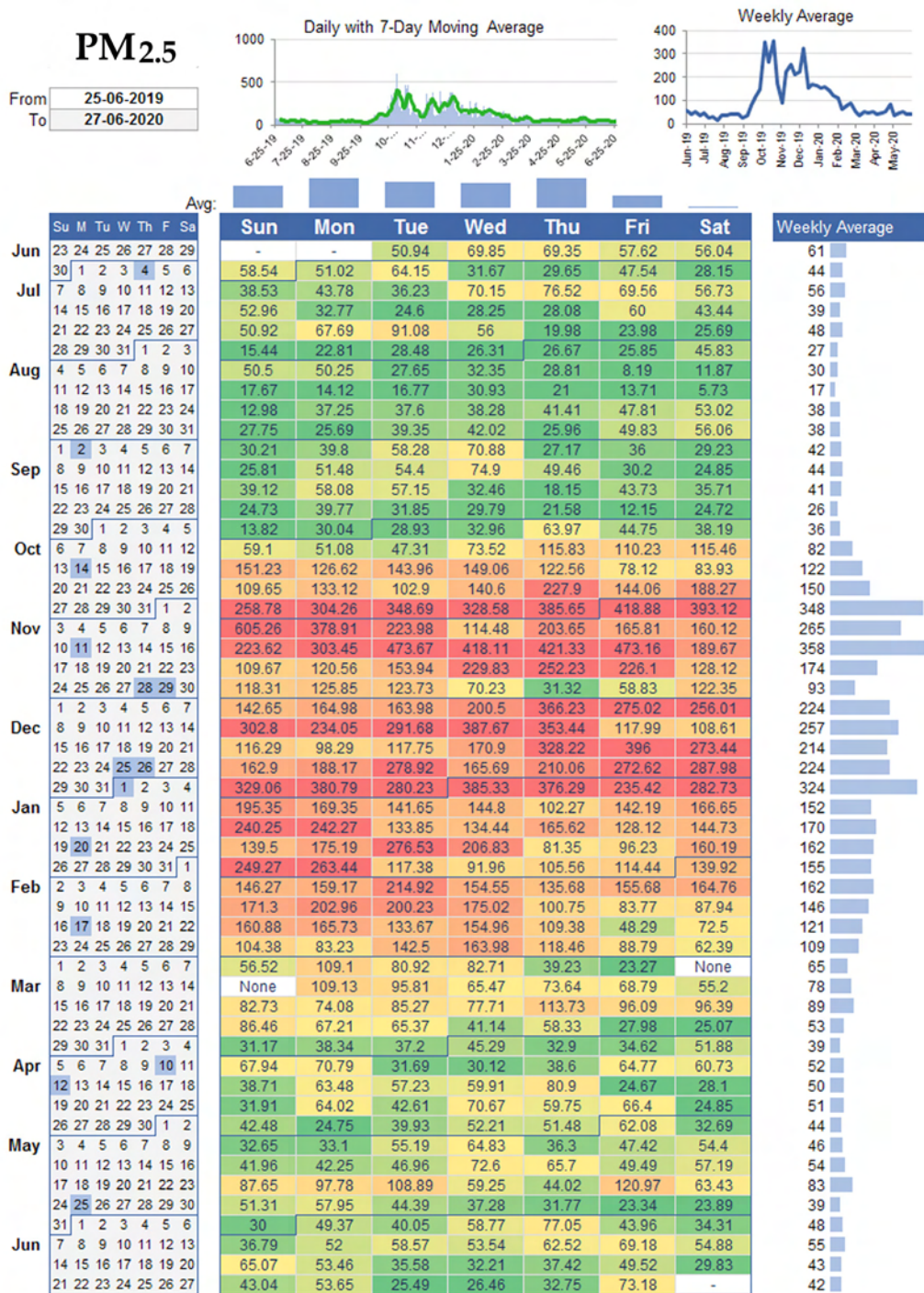
इस अध्ययन के लिए हमने दिल्ली के चयनित इलाके (अशोक विहार) के विभिन्न वायु प्रदूषकों पर डेटा एकत्र किया है। हमने 06 महत्वपूर्ण प्रदूषकों (PM2.5, PM10, NO₂, NH₃, CO, एवं CO₂) का डाटा केंद्रीय प्रदूषण नियंत्रण बोर्ड के ऑनलाइन सिस्टम के माध्यम से डाउनलोड किया (<https://app-cpcbcr-com/ccr/#/caaqm-dashboard-all/caaqm-landing>)। हमने एक वर्ष के लिए PM2.5 और PM10 की सांद्रता (Concentration) के डाटा का उपयोग इस अनुसन्धान पत्र में किया है। ऐसा करने के पीछे के कारणों में महत्वपूर्ण रूप से हमने इन कणों के प्रभाव को स्वास्थ्य के खतरे के रूप में उजागर करने और कोविड-19 प्रेरित राष्ट्रव्यापी लॉकडाउन के कारण परिवेशी वायु में इन कणों की कमी को दर्शाने के लिए किया है। एकत्र किए गए सभी डेटा को एक्सेल-2010 का उपयोग करके सरल सारणीयन विधि से सांख्यिकी रूप से विश्लेषण किया गया है। हमने इस अनुसन्धान पत्र में विभिन्न वैज्ञानिक प्रकाशन, वैज्ञानिक रिपोर्ट, विशेषज्ञों की राय, विशेषज्ञों के साक्षात्कार, स्थानीय और राष्ट्रीय समाचार पत्रों के माध्यम से पूरे देश में हवा की गुणवत्ता पर लॉकडाउन के प्रभाव की रिपोर्ट का भी अध्ययन किया है। प्रस्तुत शोधपत्र में हमने भारत की राजधानी नयी दिल्ली से चुने हुए एक क्षेत्र अशोक विहार के पिछले एक वर्ष के दौरान (25-06-2019 से 27-06-2020) की हवा की गुणवत्ता का अध्ययन किया। हमने चुने गए क्षेत्र की प्रतिदिन की हवा के नमूनों का विश्लेषण करते हुए वायु प्रदूषण में शामिल कणिका तत्वों PM2.5 एवं PM10 के स्तर को नापा तथा साथ ही साथ वायु में इन कणिका तत्वों के स्तर पर लॉकडाउन से पहले, लॉकडाउन के दौरान एवं लॉकडाउन के बाद (अनलॉक -1) के प्रभावों का अध्ययन किया। PM2.5 और PM10 के स्तर को हमने राष्ट्रीय वायु गुणवत्ता सूचकांक (National Air Quality Index, NAQI) के आधार पर विश्लेषित किया है जिसमें विभिन्न रंगों द्वारा हवा की गुणवत्ता के पैटर्न को सहज रूप से प्रदर्शित किया गया है।

चयनित क्षेत्र (अशोक विहार, नयी दिल्ली) में पिछले एक साल के वायु के नमूने में कणिका तत्वों की सांद्रता का विश्लेषण

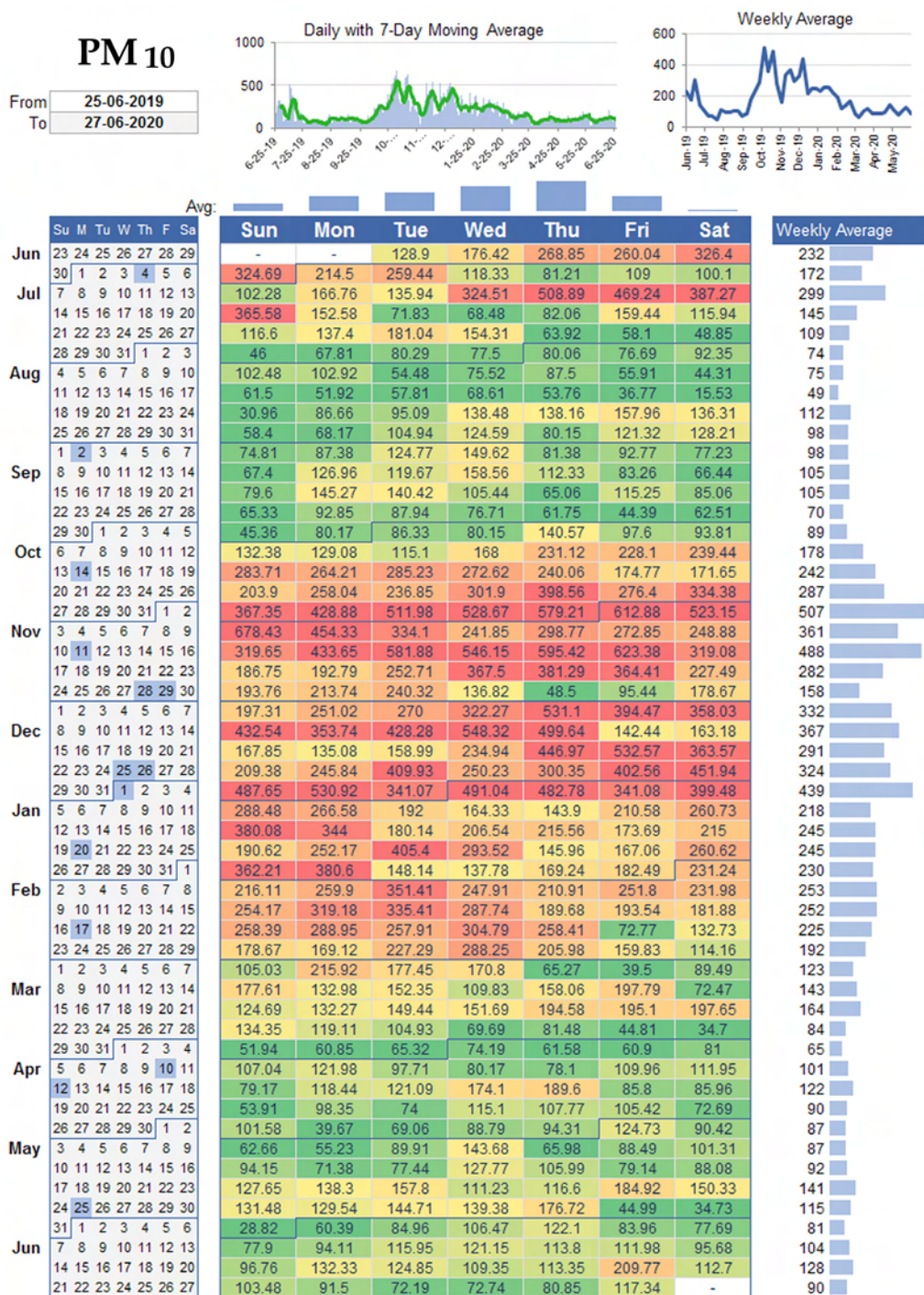
अनुसंधान के परिणाम तालिका -01 एवं तालिका -02 के द्वारा प्रस्तुत किये गए हैं। तालिका-01 के अवलोकन से स्पष्ट है कि लॉकडाउन के कारण PM2.5 की सांद्रता के साप्ताहिक औसत में कमी आयी है जिसकी वजह से हवा की गुणवत्ता में सुधार परिलक्षित हुआ है। PM2.5 की सांद्रता में लॉकडाउन के प्रथम सप्ताह में ही 40.44% की कमी दर्ज की गयी जो निरंतर कम होती गयी। जबकि लॉकडाउन के उपरान्त, अनलॉक -01 के प्रथम सप्ताह में ही PM2.5 की सांद्रता में 14.58% का उछाल देखा गया। प्रस्तुत आकड़ों के अवलोकन से यह प्रमाणित होता है कि लॉकडाउन का वायु की गुणवत्ता के सुधार में योगदान है (तालिका -01)। साथ ही, तालिका -02 यह स्पष्ट दर्शाती है कि लॉकडाउन के प्रथम सप्ताह में PM10 की सांद्रता में 48.78% की जबरदस्त कमी दर्ज की गयी। PM10 सांद्रता में यह कमी लॉकडाउन के अन्तिम सप्ताह तक लगातार दर्ज की गयी, ये गिरावट अनलॉक के साथ थम गयी और लॉकडाउन के बाद PM10 सांद्रता में वृद्धि होना फिर से शुरू हो गया। प्राप्त परिणामों से यह प्रदर्शित हुआ है कि लॉकडाउन के दौरान चयनित प्रदूषकों PM2.5 और PM10 की सांद्रता में काफी कमी हुई। अशोक विहार निगरानी स्टेशन पर प्रमुख वायु प्रदूषकों जैसे PM2.5, PM10, NO₂, NH₃, CO, एवं CO₂ की सांद्रता में लॉकडाउन से पहले और लॉकडाउन के दौरान (01 मार्च 2020 से 21 मार्च 2020 तथा 25 मार्च 2020 से 31 मई 2020 के बीच दैनिक औसत के आधार पर 35.39, 31.55, 54.23, 0.47, 92.0 एवं 10.34% की कमी देखी गयी।

वर्ष 2019 एवं वर्ष 2020 के दौरान अशोक विहार निगरानी स्टेशन पर प्रमुख प्रदूषकों जैसे PM2.5, PM10, NO₂, NH₃, CO, एवं CO₂ की मात्रा में (25 मार्च 2019 से 31 मई 2019 तथा 25 मार्च 2020 से 31 मई 2020) 47.37, 56.16, 67.34, 18.05, 56.92 एवं 38.3% की कमी दर्ज की गयी है। यह अध्ययन दर्शाता है कि यदि प्रदूषण स्रोत पर नियंत्रण कर लिया जाये तो वायु की गुणवत्ता को बेहतर किया जा सकता है। साथ ही, लॉकडाउन जैसे अस्थायी विकल्प के कार्यान्वयन से वायु की गुणवत्ता में काफी सुधार हो सकता है और भविष्य में इसे वायु प्रदूषण में कमी करने के लिए वैकल्पिक उपाय के रूप में व्यवहार में लाया जा सकता है [5]। तालिकाओं के अध्ययन से यह भी प्रदर्शित हुआ कि प्रदूषकों की सांद्रता भी विषयगत परिस्थितियों में अंतर-मौसमी असमानता के साथ उतार-चढ़ाव कर सकती है। उदाहरण के लिए, भारतीय उपमहाद्वीप में मानसून के महीनों (जून से सितंबर) में PM2.5 और PM10 की सांद्रता साल के अन्य महीनों की तुलना में बहुत कम रहती है [6]। तालिकाओं में राष्ट्रीय वायु सूचकांक-2014 के अनुसार परिवेशी वायु में मौजूद प्रदूषकों की मात्रा का स्वास्थ्य पर प्रभाव को भिन्न रंगों के द्वारा दर्शाया गया है जिसमें 0 से 30 तक अच्छा (हरा), 31 से 60 तक माध्यम (पीला), 61 से 90 तक संवेदनशील समूह के लिए अस्वास्थ्यकर (नारंगी), 91 से 120 तक अस्वास्थ्यकर (भूरा), 121 से 250 तक बहुत ही अस्वास्थ्यकर (नीला) एवं 250 से अधिक को खतरनाक (लाल) माना जाता है।

डॉ० शिल्पी शाक्य एवं डॉ० बिन्ध्याचल यादव, "वायुमंडलीय कणिकीय द्रव्यों PM2.5 एवं PM10 के स्तरों पर लॉकडाउन"



तालिका-01: - नयी दिल्ली के अशोक विहार में जून 2019 से जून 2020 की अवधि के दौरान दैनिक (24 घंटे) औसत PM2.5 सांद्रता की स्थिति।



तालिका-02: - नयी दिल्ली के अशोक विहार में जून 2019 से जून 2020 की अवधि के दौरान दैनिक (24 घंटे) औसत PM₁₀ सांद्रता की स्थिति।

निष्कर्ष

सम्पूर्ण विश्व सहित भारत में वायु प्रदूषण का प्रभाव, मौतों और बेहतर जीवन की प्रत्याशा पर पड़ता है [7]। कई अध्ययनों से यह बात साफ हो गयी है कि उच्च आय वाले देशों की तुलना में कम आय और मध्यम आय वाले देशों में वायु प्रदूषण का अधिक प्रभाव है [8]। विश्व स्तर पर भारत में वायु प्रदूषण की उच्चतम मात्रा एक खतरनाक एवं चिंतनीय विषय है। वायु प्रदूषण के प्रमुख घटक हमारे परिवेश में मौजूद परिवेशी कण पदार्थ का प्रदूषण, घरेलू वायु प्रदूषण और वायुमंडल की सबसे निचली परत में मौजूद ओजोन प्रदूषण हैं [9]। भारत में बिजली उत्पादन हेतु उपयोग किये जाने वाले कोयले, उद्योगों से निकला उत्सर्जन, निर्माण गतिविधियाँ, ईंट भट्टों का धुआँ, परिवहन, सड़कों की धूल इत्यादि प्रमुख हैं। भारत में परिवेशी कणों (PM2.5) का स्तर दुनिया का सबसे अधिक है। यहाँ यह महत्वपूर्ण तथ्य है कि मानव स्वास्थ्य को लाभ पहुँचाने के अलावा, भारत में वायु प्रदूषण पर नियंत्रण का पारिस्थितिकी तंत्र के कई पहलुओं पर व्यापक लाभकारी प्रभाव, जिसमें पशु और पौधों का स्वास्थ्य भी शामिल है, परिलक्षित होगा। COVID-19 सर्वव्यापी महामारी के खिलाफ एक निवारक उपाय के रूप में, केंद्र सरकार ने 25 मार्च, 2020 को देशव्यापी लॉकडाउन का आदेश दिया। देश में लॉकडाउन को कई बार प्रतिबंधों में क्रमिक छूट के साथ बढ़ाया गया है। उपग्रह से प्राप्त डाटा ने देश के अधिकांश हिस्सों में COVID-19 लॉकडाउन के बाद पार्टिकुलेट मैटर या एयरोसोल के स्तर में महत्वपूर्ण गिरावट दिखाई है [10]। लॉकडाउन के परिणामस्वरूप भारत में विशेषज्ञों ने ऐसे कई पर्यावरणीय कारकों को रेखांकित किया है जो कम औद्योगिक और मानवीय गतिविधियों के कारण, हवा की गुणवत्ता, ध्वनि प्रदूषण, जल की गुणवत्ता और जैव विविधता में हुए आश्चर्यजनक सुधार के रूप में सामने आए हैं [11]। नयी दिल्ली को अंतर्राष्ट्रीय स्तर पर अपने चरम वायु प्रदूषण स्तर के लिए जाना जाता है। लॉकडाउन में चयनित प्रदूषक PM2.5 और

PM10 में सर्वाधिक कमी देखी गई। वर्तमान लेख में राष्ट्रीय राजधानी शहर दिल्ली की वायु गुणवत्ता पर भारत में COVID-19 सर्वव्यापी महामारी के तेजी से प्रसार को रोकने के लिए लगाए गए लॉकडाउन के प्रभाव का मूल्यांकन, राष्ट्रीय वायु गुणवत्ता सूचकांक (NAQI) और प्रमुख प्रदूषकों की सांद्रता के आधार पर किया गया है। भारत में वायु गुणवत्ता को स्थानीय कारक भी प्रभावित करते हैं [12]। भारत में नवंबर के बाद से फरवरी तक सर्दियों का मौसम और उसके बाद नियमित तापमान में बढ़ोत्तरी होती है। इस दौरान स्थानीय प्रदूषकों में वृद्धि हो जाती है और इस मौसम के दौरान प्रदूषण स्तर में वृद्धि होती है। इसलिए, लॉकडाउन को वैकल्पिक नीति उपाय के रूप में मानने के लिए, क्षेत्रीय मौसम संबंधी स्थितियों के साथ प्रदूषक सांद्रता के मौसमी बदलाव की जाँच भी आवश्यक है। हमने यह भी देखा है कि राष्ट्रीय वायु गुणवत्ता सूचकांक में काफी कमी पूरी दिल्ली में लॉकडाउन की अवधि के दौरान देखी गई है। लॉकडाउन के शुरू होने के बाद हवा की गुणवत्ता में उल्लेखनीय सुधार हुआ। लॉकडाउन में कुछ दिन में ही PM2.5 और PM10 की सांद्रता अनुमन्य सीमा के भीतर आ गई।

वर्तमान में अनलॉक - 3 (01 अगस्त, 2020 से 31 अगस्त, 2020 तक) एवं अनलॉक - 4 (01 सितम्बर, 2020 से अद्यतन प्रभावी) में वायु प्रदूषकों के रूप में PM10 की मात्रा में प्रतिदिन दैनिक औसत के आधार पर 105.53% की वृद्धि दर्ज की गयी एवं PM 2.5 की मात्रा में 103% की वृद्धि दर्ज की गयी। उक्त अवधि के दौरान अमोनिया प्रदूषक के रूप में दैनिक औसत के आधार पर 9.3% की कमी दर्ज की गयी है (अमोनिया में यह कमी स्थानीय वायुमंडलीय कारकों तथा मानसून की सक्रियता के कारण भी हो सकती है)। अतः स्पष्ट है कि अनलॉक के दौरान वायु प्रदूषकों की मात्रा में इस प्रकार की वृद्धि लॉकडाउन के दौरान किये गए उपायों को अधिक तर्कसंगत बनाती है।

भारत के लिए स्वच्छ वायु योजना: स्वच्छ पर्यावरण के लिए नई दिशा

भारतीय शहरों में वायु प्रदूषण एक गंभीर मुद्दा है। विश्व स्वास्थ्य संगठन की वर्ष 2018 की रिपोर्ट के अनुसार दुनिया के सबसे प्रदूषित 15 शहरों में से 14 भारत के हैं। भारत के कई हिस्सों में वायु प्रदूषण का उच्च स्तर मौजूद है इसलिए कई शहर वायु प्रदूषण संकट को दूर करने के लिए प्रयास कर रहे हैं। इन कार्यों में वायु गुणवत्ता निगरानी नेटवर्क को मजबूत करना, परिवहन क्षेत्र से प्रदूषण को कम करने के लिए इलेक्ट्रिक वाहनों को अपनाना, नियामक अनुपालन को मजबूत करना और औद्योगिक उत्सर्जन को बेहतर ढंग से नियंत्रित करना इत्यादि शामिल है। यद्यपि देश में वायु प्रदूषण को कम करने हेतु कई योजनाओं पर कार्य हो रहा है एवं कई योजनाओं में अपेक्षित प्रगति भी दर्ज की गयी है। भारत के शहरों के लिए स्वच्छ वायु योजनाओं का कार्यान्वयन केवल तभी काम कर सकता है जब स्पष्ट लक्ष्य और जवाबदेही तथा सम्बन्धित सरकारी एजेंसियों के बीच पर्याप्त समन्वय हो। विभिन्न एजेंसियों के कार्यों में समन्वय की कमी और किसी एक की जवाबदेही न होने के कारण कई योजनाएं धरातल पर उतर ही नहीं पाती हैं।

वायु प्रदूषण रोकथाम के लिए दिल्ली सरकार द्वारा उठाए गए कदम

भारत की राजधानी दिल्ली पर्यावरण क्षतिपूर्ति शुल्क लागू करने वाला पहला राज्य बन गया है। इस योजना के तहत सभी नए वाहनों को पर्यावरण क्षतिपूर्ति शुल्क के लिए न्यूनतम राशि का भुगतान करना होगा। इसके साथ ही, दिल्ली सरकार द्वारा बड़े पैमाने पर ई-रिक्शा के सबसे बड़े नेटवर्क की अनुमति दी है। ऐसे मौसम, जिनमें वायु प्रदूषण अपने चरम पर होता है, उस दौरान निर्माण कार्यों हेतु भारी जुर्माना लगाये जाने की व्यवस्था की गयी है। दिल्ली, कोयले पर पूर्ण प्रतिबंध लगाने वाला पहला राज्य है। यहाँ हरियाली कवर की व्यापक वृद्धि के निरंतर प्रयास किये जा रहे हैं।

भारत में वायु प्रदूषण के सम्बन्ध में वायु गुणवत्ता पर केंद्रीकृत डेटाबेस का अभाव है जिसके कारण लंबी अवधि के अध्ययन के लिए डेटा की उपलब्धता एवं विश्लेषण करना मुश्किल है। राष्ट्रीय वायु गुणवत्ता निगरानी कार्यक्रम स्टेशनों से दीर्घकालिक डेटासेट की पहुंच और उपलब्धता जैसे हालिया प्रयास एक स्वागत योग्य कदम है, जो निस्संदेह इस विषय पर काम करने की सुविधा प्रदान करते हैं। पिछले दो दशकों में भारत में वायु गुणवत्ता निगरानी के संबंध में महत्वपूर्ण सुधार किए गए हैं, और प्रमुख शहरों में निरंतर परिवेशी वायु गुणवत्ता निगरानी स्टेशनों की स्थापना पर ध्यान केंद्रित करने के साथ राष्ट्रीय निगरानी कार्यक्रम को और मजबूत करने का प्रयास चल रहा है।

सुनहरे भविष्य की दिशा एवं सतत विकास के लिए अभिनव समाधान

वायु प्रदूषण की गंभीर समस्या और इसके जुड़े स्वास्थ्य संबंधी खतरे कई वैज्ञानिक लेखों के माध्यम से अब अच्छी तरह से स्थापित हैं [13]। देश में सबसे बड़ी वैश्विक पर्यावरण चुनौतियों में से वायु प्रदूषण के समाधान के लिए पर्यावरण वन और जलवायु परिवर्तन मंत्रालय द्वारा राष्ट्रीय स्वच्छ वायु कार्यक्रम (National Clean Air Programme) शुरू किया गया। वायु प्रदूषण पर अंकुश लगाने के लिए भारत के दृढ़ संकल्प और स्वच्छ वायु कार्यक्रम का समर्थन करने के उद्देश्य से वर्ष 2024 तक PM2.5 और PM10 की सांद्रता को 20-30 प्रतिशत तक कम करने का लक्ष्य है [14]। इस कार्यक्रम के तहत चार भारतीय शहरों में एक स्वच्छ वायु परियोजना शुरू की गई है। भारत में द एनर्जी एंड रिसोर्स इंस्टीट्यूट (The Energy and Resources Institute, TERI) ने स्विस् एजेंसी फॉर डेवलपमेंट एंड कोऑपरेशन (Swiss Agency for Development and Co-operation, SDC) के साथ मिलकर वायु गुणवत्ता में सुधार के लिए भारत में पहले स्वच्छ वायु कार्यक्रम की घोषणा की है।

यह परियोजना उत्तर प्रदेश के लखनऊ और कानपुर तथा महाराष्ट्र के पुणे और नासिक में शुरू की गई है। कोविड-19 लॉकडाउन ने चार प्रमुख शहरों में भारत के स्वच्छ वायु कार्यक्रम द्वारा निर्धारित लक्ष्य का 95% हासिल किया है। [15] लॉकडाउन के दौरान प्रदूषण के स्तर में भारी गिरावट भारत के वायु प्रदूषण प्रबंधन में सबक देती है जिसको देश में स्वच्छ वायु लक्ष्यों को प्राप्त करने में शामिल करने की आवश्यकता है। वायु प्रदूषण पर्यावरण के लिए सबसे बड़ा खतरा है और सभी को प्रभावित करता है। वायु प्रदूषण विषाक्त पदार्थों के वातावरण में उपस्थिति के कारण होता है, जो मुख्य रूप से मानव गतिविधियों द्वारा उत्पन्न होता है। यह पर्यावरण के लिए सबसे बड़ा खतरा है और मनुष्य के साथ साथ जानवर, फसल, शहर, वन, जलीय पारिस्थितिक तंत्र सब को ही प्रभावित करता है। कभी-कभी यह प्राकृतिक घटनाओं जैसे ज्वालामुखी विस्फोट, धूल भरी आंधी और जंगल की आग से भी हो सकता है, जिससे वायु की गुणवत्ता में भी कमी आती है। वायु प्रदूषण मुख्य रूप से बिजली, परिवहन, उद्योग, आवासीय, निर्माण और कृषि सहित कई क्षेत्रों से उत्पन्न होता है। भारत में वायु प्रदूषण की समस्या से निपटने के लिए कई अभिनव उपाय सुझाए गए हैं [16]। इनमें से कुछ उपायों में (1) मौसम पौधे की दीवारें खड़ी करना (2) ऑक्सीजन चेम्बर का निर्माण, जिनमें लोग आकर प्रदूषण मुक्त हवा में साँस ले सकते हैं (3) इलेक्ट्रिक कारों का प्रयोग (4) वर्टिकल फॉरेस्ट सिटी (5) विशाल स्प्रिंकलर आदि शामिल हैं। सरकार द्वारा लगातार इन उपायों के बारे में नीति एवं योजनाएँ बनाई जाती हैं। कोरोना संक्रमण से हमारे सामने कई पहलू उजागर हुए हैं, उन पहलुओं में से एक पहलू यह भी है कि पूरे विश्व के साथ साथ भारत की हवा भी लॉकडाउन की वजह से स्वच्छ हो गयी है। ग्रीन हाउस गैसों की मात्रा में भी कमी के संकेत मिले हैं। इन सबके साथ कई नकारात्मक चीजें भी देखने को मिल रही हैं— हमारे द्वारा उपयोग में लाई जा रही प्लास्टिक की मात्रा में कई गुना वृद्धि हो

गयी है, साथ ही साथ सड़को पर निजी वाहनों की संख्या में लगातार बढ़ोत्तरी हो रही है। इन सबका हमारे पर्यावरण पर ऋणात्मक प्रभाव पड़ना तय है [17]। समय की मांग ऐसे रचनात्मक उत्तर की है, जो कई चुनौतियों से एक साथ निपट सके। वायु प्रदूषण से निपटने हेतु कई सुझाव हैं पर इन सब को अमल में लाना अपने आप में एक बड़ी चुनौती है। हमारा कल और बेहतर तब ही हो सकेगा जब हम सब मिलकर इस दिशा में कदम से कदम मिला कर चलेंगे। सबसे अहम सवाल यह उठता है कि क्या कोरोना के बाद जीवन पहले जैसा हो जाएगा? हमें यह लगता है कि इसका उत्तर अभी "नहीं" है। COVID-19 के बाद जीवन सामान्य हो जाने की उम्मीद तो है परन्तु यह बात भी निश्चित है कि COVID-19 के बाद का हमारा जीवन पहले जैसा नहीं होगा। इसलिए, हमारे जीवन पर COVID-19 के प्रभाव के बारे में विभिन्न क्षेत्रों में अधिक शोध की आवश्यकता है। हमारे दैनिक जीवन पर सर्वव्यापी महामारी के प्रभाव को समझना, अनुसंधान हेतु नए क्षेत्रों का निर्माण करेगा जिससे हमें न केवल COVID-19 के नकारात्मक प्रभावों को कम करने में मदद मिलेगी बल्कि सक्रिय समाधान तैयार करने में मदद मिलेगी जो COVID-19 समस्या के उपरांत हमारे जीवन को आकार देगा।

प्रतिस्पर्धी वित्तीय हित की नकार घोषणा

वर्तमान अध्ययन के लेखक यह घोषणा करते हैं कि उनके ऐसे कोई ज्ञात प्रतिस्पर्धी वित्तीय हित, रुचियाँ या व्यक्तिगत संबंध नहीं हैं जो इस अध्ययन में बताए गए काम को प्रभावित कर सकते थे।

आभार

इस अनुसन्धान पत्र के लेखक इस अध्ययन के दौरान दिए गए सुझावों, महत्वपूर्ण और विश्लेषणात्मक टिप्पणियों एवं अनुसन्धान पत्र में उल्लिखित तालिकाओं हेतु अपने सहयोगी सुशांत महतो, भूगोल विभाग गौर बंग विश्वविद्यालय, पश्चिम बंगाल का आभार व्यक्त करते हैं।

सन्दर्भ सूची

1. Agarwal, Aviral, et al. Comparative study on air quality status in Indian and Chinese cities before and during the COVID-19 lockdown period. *Air Quality, Atmosphere & Health*. 2020: 1-12.
2. Al-Kindi, Sadeer G., et al. Environmental determinants of cardiovascular disease: lessons learned from air pollution. *Nature Reviews Cardiology*. 2020: 1-17.
3. Arulprakasajothi, M., et al. An analysis of the implications of air pollutants in Chennai. *International Journal of Ambient Energy*. 2020: 209-213.
4. Combes, Alain, and Guillaume Franchineau. Fine particle environmental pollution and cardiovascular diseases. *Metabolism*. 2019: 153944.
5. Gautam, Sneha. The influence of COVID-19 on air quality in India: a boon or inutile. *Bulletin of Environmental Contamination and Toxicology*. 2020: 1.
6. Ghosh, Sasanka, et al. Impact of COVID-19 Induced Lockdown on Environmental Quality in Four Indian Megacities Using Landsat 8 OLI and TIRS-Derived Data and Mamdani Fuzzy Logic Modelling Approach. *Sustainability*. 2020: 5464.
7. INDIA, POMPI. Census of India 2011 Provisional Population Totals. New Delhi: Office of the Registrar General and Census Commissioner. 2011.
8. Isaifan, R. J. The dramatic impact of Coronavirus outbreak on air quality: Has it saved as much as it has killed so far?. *Global Journal of Environmental Science and Management*. 2020: 275-288.
9. कुमार आशुतोष, नाइट्रिक ऑक्साइड : एक संक्षिप्त चर्चा, विज्ञान प्रकाश, 2010, 8, 1-4 , पृष्ठ संख्या 16-18.
10. Mahato, Susanta, Swades Pal, and Krishna Gopal Ghosh. Effect of lockdown amid COVID-19 pandemic on air quality of the megacity Delhi, India. *Science of the Total Environment*. 2020: 139086.
11. Maji, Sanjoy, Santu Ghosh, and Sirajuddin Ahmed. Association of air quality with respiratory and cardiovascular morbidity rate in Delhi, India. *International journal of environmental health research*. 2018: 471-490.
12. Rajak, Rahul, and Aparajita Chattopadhyay. Short and long-term exposure to ambient air pollution and impact on health in India: a systematic review. *International journal of environmental health research*. 2019: 1-25.
13. Saadat, Saeida, Deepak Rawtani, and Chaudhery Mustansar Hussain. Environmental perspective of COVID-19. *Science of the Total Environment*. 2020: 138870.
14. Schraufnagel, Dean E., et al. Air pollution and noncommunicable diseases: A review by the Forum of International Respiratory Societies Environmental Committee, Part 2: Air pollution and organ systems. *Chest*. 2019: 417-426.
15. Shukla, Komal, et al. Mapping spatial distribution of particulate matter using Kriging and Inverse Distance Weighting at supersites of megacity Delhi. *Sustainable Cities and Society* 2020: 101997.
16. Tiwari, S., et al. Assessments of PM1, PM2.5 and PM10 concentrations in Delhi at different mean cycles. *Geofizika*. 2012: 125-141.
17. Zou, Xiang, et al. Environment and air pollution like gun and bullet for low-income countries: war for better health and wealth. *Environmental Science and Pollution Research*. 2016: 3641-3657.

भारतीय उद्योग के पितामह – जमशेदजी टाटा

Father of Indian Industry – Jamshedji Tata

प्रो. वीरेन्द्र ग्रोवर,
प्रबन्ध सम्पादक “उद्योग संचेतना”

vngrover@gmail.com

संक्षिप्त परिचय :

पारसी पुजारी परिवार में 3 मार्च, 1839 को जन्मे जमशेदजी नुसरवानजी टाटा आज उद्यमिता के अवतार हैं। उन्होंने 1868 में टेक्सटाइल मिल से जिस उद्योग समूह की नींव रखी थी, उसके आज 10 वर्गों (वर्टिकल्स) में तीस उद्योग हैं। उद्योग निर्माण में श्रमिक कल्याण की जो पहल एम्प्रेस मिल से हुई उसकी परिणति जमशेदपुर के स्टील संयंत्र की टाउनशिप से होती है। स्वदेशी वस्त्र के बाद स्टील संयंत्र, जलविद्युत संयंत्र और राष्ट्रीय प्रतिभा पोषण के लिए उच्च शिक्षण संस्थान की परिकल्पना को लेकर जमशेदजी ने 1880 से 1904 तक संघर्ष किया, जो उनके प्रयाण के बाद ही मूर्त रूप ले सके, लेकिन जिस मेक-इन-इंडिया की भावना का बीजारोपण जमशेदजी ने किया, वह विशाल वट वृक्ष आज के युवाओं तथा उद्यमियों के लिए प्रेरणा है।

Brief Introduction :

Born in a family of Parsi priests on March 3, 1839, Jamshedji Nusserwanji Tata is an incarnation of entrepreneurship. The story that started with a Textile Mill in 1868, has turned into a global enterprise with 30 sub-enterprises in 10 verticals. Though, the 3 big dreams he saw, couldn't materialise till his death in 1904, but the spirit he lived is alive even today inspiring both youth and the entrepreneurs alike.

समग्रतामय उद्यम :

जमशेदजी टाटा उद्यमिता की अनूठी मिसाल थे। राष्ट्र के लिए पूँजी निर्माण, लोककल्याण, जनसशक्तीकरण, तथा स्वावलंबन, उनके उद्यम के मूल सिद्धान्त थे। टाटा समूह में कुछ बातें शुरू से अनन्य रूप से जुड़ी रही हैं, जैसे कि उद्यम में प्रौद्योगिकी का समुन्नयन एवं नवाचारों का समावेश, मानव संसाधन का सतत् उन्नयन, तथा परस्पर सहयोग तथा सहकारिता। टाटा ग्रुप ने अपने उद्यम में तकनीकी सुधार तथा नवाचार को हमेशा स्थान दिया। उनकी कोशिश रही कि समय के साथ संगठन के कर्मचारीगण नये कौशल सीखें तथा बदलते समय में प्रतिस्पर्धी बने रहें। टाटा समूह में इस बात पर भी हमेशा बल दिया गया कि विभिन्न इकाइयों के मध्य परस्पर सहयोग तथा सहकारिता की भावना का विकास हो। टाटा ग्रुप के प्रबन्धन में कार्मिकों का कल्याण सर्वोपरि रहा है। राष्ट्र के लिए संपदा का सृजन, कार्मिकों का कल्याण तथा उनमें संगठन के प्रति आत्मिक जुड़ाव, टाटा समूह के मंत्र रहे हैं। टाटा समूह के उत्पाद तथा सेवाओं का हेतु मानव कल्याण रहा है। इसलिए टाटा ग्रुप में प्रबन्धन तथा कार्मिकों के मध्य परस्पर आत्मीयता का सम्बन्ध रहा है। यह एक विरासत है जिसका नैरन्तर्य आज भी कायम है। यह बात स्वयं में अनुपम तथा अतुलनीय है। प्रस्तुत है टाटा समूह की नींव रखने वाले उद्योगपति माननीय जमशेदजी टाटा का संक्षिप्त जीवन परिचय जो बेहद रोचक, पठनीय एवं मननीय ही नहीं, अनुकरणीय भी है।



जमशेदजी टाटा

टाटा उद्योग समूह के जनक और पितामह जमशेदजी टाटा (1839–1904) केवल उद्यमी नहीं, उद्यम अवतार थे, जिन्होंने देश में औद्योगीकरण की नींव ऐसे समय में रखी, जब विदेशी शासक यहाँ विकास नहीं चाहते थे। नवसारी (गुजरात) के साधारण पारसी पुजारी परिवार में 3 मार्च 1839 को जन्मे, माता-पिता की पहली व इकलौती संतान में ऐसा कोई संकेत नहीं था कि एक दिन वह अपना भाग्य निर्माता बनेगा। जिस परिवार में पीढ़ियों से लोग पुजारी बनते आये थे, उस परम्परा से हट कर व्यापार में भाग्य आजमाने वाला वह पहला सदस्य था।

प्रारम्भिक शिक्षा नवसारी में लेकर 14 वर्ष की आयु में बम्बई के एलफिंस्टन कॉलेज से 1858 में “ग्रीन स्कॉलर” उपाधि ग्रहण की जो आज के स्नातक के समतुल्य है। इस शिक्षा का असर यह हुआ कि जमशेदजी जिंदगी भर के लिए शैक्षणिक व्यवस्था व पुस्तक पढ़ने के प्रेमी बन गए। लेकिन उनका यह शौक कुछ ही समय में पीछे छूट गया जब उन्हें लगा कि जीवन का लक्ष्य कुछ और ही है – बिजनेस।

जे. आर. डी. टाटा के शब्दों में, यह एक ऐसा संक्रमण काल था जब किसी स्थानीय युवक के लिए व्यापार की बात सोचना भी मुश्किल था। सन् 1857

की क्रांति के दो साल ही बीते थे और विदेशी शासकों का भय चरम पर था। बीस वर्ष की उम्र में जमशेदजी ने पिता के छोटे से धन लेन-देन के कारोबार में हाथ बँटाना शुरू किया और 9 वर्ष के अनुभव के बाद 21 हजार रुपये की पूँजी से ट्रेडिंग कम्पनी शुरू की।

इस दौरान इंग्लैंड की यात्रा में जमशेदजी ने टेक्सटाइल व्यापार को समझा और महसूस किया कि देश में ब्रिटिश प्रभुत्व वाले टेक्सटाइल उद्योग में भारत के उद्यमियों को दखल देने की काफी संभावना है। मुंबई के चिंचपोकली इलाके में एक बंद पड़ी पुरानी तेल मिल खरीद कर वहाँ अलेक्जेन्द्रा मिल के नाम से कॉटन मिल शुरू की। दो साल बाद अच्छे मुनाफे पर मिल बेचकर एक बार फिर इंग्लैंड की यात्रा पर निकल कर लंकशायर कॉटन ट्रेड का बारीकी से अध्ययन किया। मशीनों, कारीगरों व उत्पाद की गुणवत्ता देखकर जमशेदजी को विश्वास था कि वे भारत में उपनिवेशवादी शासकों को टक्कर दे सकते हैं।

परिवर्तन का बीजारोपण

एक दकियानूसी सोच थी कि कोई नई परियोजना स्थापित करनी हो तो मुम्बई ही उपयुक्त है, लेकिन जमशेदजी की सोच थी कि सफलता के लिए आवश्यक है कपास उत्पादकों से नजदीकी, रेल यातायात की सुविधा, पानी व ईंधन की उपलब्धता और महाराष्ट्र में नागपुर इसके लिए उचित स्थान था। वर्ष 1874 में डेढ़ लाख की प्रारम्भिक पूँजी से सेंट्रल इंडिया स्पिनिंग, वीविंग एंड मैन्यूफैक्चरिंग कंपनी की नींव रखी गई और तीन साल बाद जनवरी 1877 के पहले दिन जब महारानी विक्टोरिया भारत की महारानी (एम्प्रेस) बनी तो जमशेदजी ने इसे एम्प्रेस मिल घोषित कर दिया।

एम्प्रेस मिल से कामगार कल्याण की एक नई पहल हुई जो उन दिनों कोई नहीं जानता था। इसके साथ ही जमशेदजी के व्यस्त जीवन की महत्वपूर्ण यात्रा शुरू होती है, जिसमें बहुत लम्बी सोच थी। वर्ष 1880 से लेकर 1904 में जीवन के आखिरी क्षण तक

उनके तीन दूरदर्शिता से भरे सपने थे— स्टील प्लांट, जलविद्युत संयंत्र और एक विश्व-स्तरीय वैज्ञानिक संस्थान की स्थापना करना। उनके जीते जी चाहे ये तीनों चीजें फलीभूत नहीं हो पाईं, लेकिन जिस संकल्प शक्ति से उन्होंने इन प्रकल्पों का बीजारोपण किया उसका वैभवशाली स्वरूप आज टाटा उद्योग समूह के रूप में प्रत्यक्ष विद्यमान है।

लौह पुरुष

नई टेक्सटाइल मशीनरी की तलाश में मैनचेस्टर यात्रा के दौरान जमशेदजी को थॉमस कार्लाइल का भाषण सुनने का अवसर मिला, जहाँ से उनके मन में विश्वस्तरीय स्टील प्लांट की परिकल्पना जगी। यह एक बहुत बड़ा प्रकल्प था। जिस औद्योगिक क्रांति ने ब्रिटेन तथा अन्य देशों में विकास की लहर पैदा की, भारत उससे अछूता रह गया था। सरकारी नीतियाँ, दुर्गम स्थानों पर खोज का प्रयास और कदम-कदम पर रोड़े, हर मोड़ पर रूकावट ही नजर आती थी।

इस्पात संयंत्र के प्रकल्प को जिन रास्तों से गुजरना पड़ा, वहाँ साधारण व्यक्ति हार मान लेता लेकिन जमशेदजी अड़िग रहे। इतना ही नहीं, बहुत से व्यंग्य भी सहने पड़े। ग्रेट इंडियन पेनिन्सुलर रेलवे के चीफ कमिश्नर सर फ्रेडरिक अपकोट ने भरोसा दिया था कि “टाटा जितना भी रेल स्टील बनाएगा, वह खप जाएगा”। लेकिन जब 1912 में जमशेदजी की मृत्यु के बाद स्टील प्लांट में पहली पिंडी (पदहवज) रोल की गई तो सर फ्रेडरिक का कहीं अता पता नहीं था। परन्तु जमशेदजी की आत्मा वहाँ जरूर थी। अपने बेटे दोराब तथा चचेरे भाई आर. डी. टाटा के साथ, जिनके अथक प्रयास ने असंभव को संभव कर दिया था।

उद्योग का निर्माण सिर्फ ईंट पत्थर से नहीं होता। कामगार कल्याण की पहल एम्प्रेस मिल से हो गई थी। जमशेदजी ने काम के घंटे कम किये, कार्यस्थल को हवादार बनवाया, भविष्य निधि व ग्रेच्युटी को लागू किया जो पश्चिमी देशों में

भी तब जरूरी नहीं थी। स्टील प्लांट में कामगारों के लिए टाउनशिप की जरूरत और उसकी रुपरेखा जमशेदजी ने 1902 में ही दोराब टाटा को लिखे पत्र में व्यक्त कर दी थी। जब स्टील प्लांट के लिए स्थान का भी चयन नहीं हुआ था – “खुली सड़कें, छायादार पेड़, बगीचों के लिए काफी जगह, फुटबाल व हॉकी के लिए खेल के मैदान होने चाहिए” – इस सोच को जहाँ मूर्त रूप मिला, उसे जमशेदपुर कहना सच्ची श्रद्धांजलि है एक राष्ट्रभक्त, भविष्यद्रष्टा और उद्योग रत्न के लिए।

प्रतिभा पोषण

देश को गरीबी से छुटकारा दिलाने के लिए बुद्धिमान मस्तिष्कों के मार्गदर्शन व पोषण की आवश्यकता को देखते हुए सन् 1892 में ही जे. एन. टाटा अक्षय निधि (एंडोमेंट फण्ड) की स्थापना हो गई थी, जिसके माध्यम से से इंग्लैंड में शिक्षा के लिए टाटा छात्रवृत्ति प्रारम्भ हुई और इसी का परिणाम था कि वर्ष 1924 तक भारत की प्रशासनिक सेवाओं में चयनित हर 5 में से 2 अभ्यर्थी टाटा स्कॉलर होते थे। इसी स्रोत से इंडियन इंस्टिट्यूट ऑफ साइंस की स्थापना का विचार उपजा, लेकिन जैसा स्टील संयंत्र के साथ हुआ, जमशेदजी की तमन्ना उनके जीवनकाल में पूर्ण



इंडियन इंस्टीट्यूट ऑफ साइंस

न हो सकी। इसके लिए उन्होंने रुपरेखा बनाकर अपनी व्यक्तिगत पूंजी से 30 लाख रुपये देने का संकल्प कर लिया और तत्कालीन वाइसराय लार्ड कर्जन से लेकर स्वामी विवेकानंद तक से सहयोग का आग्रह किया था।

इस सन्दर्भ में स्वामी विवेकानन्द ने 1899 में लिखा था – “मुझे लगता है भारत में एक सामयिक आवश्यकता के प्रकल्प का प्रादुर्भाव हो रहा है जिसके दूरगामी परिणाम होंगें३..योजना न केवल राष्ट्रहित के परिप्रेक्ष्य में कमजोर कड़ियों पर प्रहार है बल्कि इसमें स्पष्ट दूरदर्शिता और मजबूत पकड़ भी है।” ऐसे समर्थन के बाद भी सोचना मुश्किल था कि बंगलुरु में भव्य इंडियन इंस्टिट्यूट ऑफ साइंस को आने में 12 वर्ष लग जायेंगे। आखिर सन् 1911 में संस्थान ने कार्य प्रारम्भ किया।

हाइड्रोइलेक्ट्रिक प्रकल्प में अपेक्षकृत कम अवरोध आये, किन्तु इसका लोकार्पण भी जमशेदजी के जीवन काल में न हो सका। फ्रैंक हैरिस ने जमशेदजी की आत्मकथा लिखते हुए कहा— “यह वो व्यक्ति था जिसका काम उसके जाने के बाद भी जीवन्त है और धारणा तथा परिपक्वता के बीच लकीर खींचना असंभव है। उनको स्मरण करते समय अहसास होता है कि एक मृत व्यक्ति कैसे जीवित प्राणियों में प्रेरणा का संचार कर रहा है।”

मुंबई का होटल ताज महल वह प्रकल्प है जो जमशेदजी अपने जीते जी देख सके। लेकिन इसका

विचार कैसे आया? कहते हैं, नगर के एक प्रतिष्ठित होटल में उन्हें प्रवेश से रोका गया था क्योंकि वे भारतीय थे। उनके पुत्र, मित्र और व्यापार सहयोगी सभी को इसमें संशय नजर आता था, और बहनों ने तो मजाक भी किया था कि “क्या तुम भोजनालय खोल रहे हो”। इन सब के बीच चार करोड़ इक्कीस लाख की लागत से 1903 में तैयार होटल तमाम सुख-सुविधाओं का अद्भुत नमूना था। यह मुंबई का पहला भवन था जिसमें बिजली थी और देश का पहला होटल, जहाँ अमेरिका के पंखे, जर्मनी की लिफ्ट, तुर्की के स्नानागार (टर्किश बाथ), अंग्रेज खानसामे थे।



स्वामी विवेकानन्द (दायें)
जो जमशेदजी टाटा के प्रेरणास्रोत रहे



होटल
ताज
महल

अमेरिकी लेखक लेविस एच. लफाम के शब्दों में — धन आग है, पृथ्वी, जल, वायु की ही तरह एक तत्वय मनुष्य इसका उपयोग औजार (टूल) की तरह कर सकता है या ईश्वर का अवतार मानकर इसके ईर्द—गिर्द नाचता रहे। जमशेदजी नुसरवानजी टाटा ने इस सम्पदा का उपयोग भारत तथा भारतवासियों के जीवन को संपन्न और समृद्ध बनाने में किया।

स्वदेशी के पितामह - जमशेदजी

स्वतंत्रता संग्राम में विदेशी वस्त्रों की होली जलाना आंदोलन का भाग था। किन्तु नष्ट करने की अपेक्षा निर्माण बड़ा कठिन काम है। और वह तब जब परिस्थितियाँ विपरीत हों तथा शासक विदेशी हो। वस्त्र उद्योग से स्वदेशी की नींव रखने वाले जमशेदजी द्वारा लगाया गया वट वृक्ष आज साधारण नमक से से लेकर अत्याधुनिक रक्षा उपकरणों के स्वदेशी उत्पादन के क्षेत्रों में संलग्न है। सूचना प्रौद्योगिकी, इस्पात, मोटर वाहन, टेलीकॉम व मीडिया आदि दस विविध श्रेणियों में टाटा समूह के वैश्विक

उद्योग संगठन में आज 30 उपक्रम हैं, जिनका कुल वार्षिक कारोबार 10,600 करोड़ अमरीकी डॉलर है। ब्रांड मूल्य के रूप में टाटा समूह देश में प्रथम स्थान पर है।

कोलकाता महानगर में हुगली (गंगा) नदी पर टाटा स्टील के भारतीय इस्पात से भारतीयों द्वारा अंतर्राष्ट्रीय मानकों के अनुरूप अभिकल्पित तथा भारतीय उद्योग द्वारा विनिर्मित पुल, मेक-इन-इंडिया की संकल्पना का केवल प्रतीक ही नहीं, प्रेरणा स्रोत है, जो सन् 2018 में 75 वर्ष पूरे कर चुका है। 3 फरवरी 1943 इतिहास का वह स्वर्णिम दिन था, जब द्वितीय विश्वयुद्ध के चलते बिना किसी भव्य आयोजन के हावड़ा ब्रिज का लोकार्पण किया गया था।

स्वास्थ्य सेवा

जमशेदजी ने स्वास्थ्य सेवा को भी व्यापार का अभिन्न अंग माना। इसकी झलक हमें कोविड-19 की वर्तमान वैश्विक महामारी काल में देखने को मिली है। बात 1896 की है जब मम्बई में प्लेग

महामारी का प्रकोप हुआ था। यूक्रेन के युवा डॉक्टर प्रोफेसर हाफकिन ने ग्रान्ट मेडिकल कॉलेज के बरामदे में जैसे-तैसे एक प्रयोगशाला स्थापित कर सन् 1897 में टीका विकसित किया। जमशेदजी ने न केवल इस कार्य को प्रोत्साहित किया बल्कि उन्होंने अपने स्टाफ, मित्रों व परिवार सहित स्वयं टीका लगवाकर समाज से मिथ्या भय और गलतफहमी को दूर करने का काम किया।

टाटा उद्योग समूह का एक
विहंगम स्वरूप
सौजन्य : टाटा उद्योग समूह
की वेबसाइट पर आधारित
हिंदी रूपांतर



पुस्तक समीक्षा

प्राचीन हिंदू विज्ञान : इसका संचरण और विश्व संस्कृतियों पर प्रभाव

ANCIENT HINDU SCIENCE: Its Transmission and Impact on World Cultures

Author: Alok Kumar. San Rafael, CA: Morgan & Claypool Publishers, 2019. 212 pages. Paperback \$29.95; Hardcover \$49.95

भारत में प्रकाशित पुस्तक—

ANCIENT HINDU SCIENCE: Its Transmission and Impact on World Cultures.

Author: Prof. Alok Kumar , Publisher: JAICO BOOKS - India; www.jaicoBooks.com

प्रो. रमन की पुस्तक—समीक्षा का भावानुवाद -

आधुनिक दुनिया में विज्ञान एक अंतर्राष्ट्रीय उद्यम है जो दुनिया भर के लोगों को लाभान्वित करता है; दुनिया भर के लोगों द्वारा अध्ययन किया जाता है और दुनिया भर के लोग विभिन्न उपायों से इसके प्रति अपना योगदान करते हैं। लेकिन विज्ञान के ये वैश्विक और परा-सांस्कृतिक पहलू मानवता के लंबे इतिहास में हाल ही में शामिल हुए हैं। अधिक प्राचीन समय में, वैज्ञानिक दिमाग आमतौर पर स्वतंत्र रूप से खोजबीन और श्रम-साधना करते थे और दूसरे स्थानों पर अन्य लोग क्या कर रहे थे इसके बारे में जानकारी या तो उन्हें बिलकुल भी नहीं होती थी या फिर बहुत कम होती थी। तथापि, समय-समय पर व्यापार और विजय के माध्यम से विचारों और अंतर्दृष्टियों का छुटपुट आदान-प्रदान और पारस्परिक मेल मिलाप होता रहता था।

सत्रहवीं शताब्दी में, पश्चिमी यूरोप में, आधुनिक विज्ञान और इसकी एकदम नई कार्यप्रणाली के उद्भव के साथ, जो प्रचुर मात्रा में फल देने वाली थी, लोग यह सोचने लगे कि प्राचीन जगत में न तो विज्ञान था और न ही गणित। इस अज्ञानता ने एक अहंकार का रूप ले लिया जिसके कारण सभी प्राचीन समाजों को वैज्ञानिक धारणाओं की बंजर भूमि मान लिया गया।

लेकिन प्राचीन ग्रीक विचारकों के लेखन की खोज ने अतीत के बारे में जानकारी के लिए एक और गंभीर खोज की मुहिम का प्रवर्तन किया। उन्नीसवीं सदी में बेबीलोन और मिस्र का सम्पन्न विज्ञान प्रकाश में आया, और जल्द ही चीन और भारत के कई वैज्ञानिक खजानों पर से भी पर्दा हटा। समान रुचि का विषय यह था कि विज्ञान के इन कोषों का आदान-प्रदान कैसे हुआ था और क्यों उन में से कुछ की चमक अपने मूल स्थान पर भी विलुप्त हो गई थी।

प्राचीन सभ्यताओं के विज्ञानों के बहुत सारे टुकड़ों को जोड़ कर एक ज्ञान निकाय के रूप में प्रस्तुत करने वाले अनगिनत विद्वानों और इतिहासकारों का हम पर बहुत बड़ा ऋण है। इससे उन विज्ञानों का एक विशाल निकाय निर्मित हुआ जो दुनिया के विभिन्न हिस्सों में प्रस्फुटित हुए थे। प्राचीन रचनात्मकता संबंधी अल्पज्ञात तथ्यों और विकृत दृष्टि पर नई रोशनी डालने के प्रयास जारी हैं।

लेकिन प्राचीन विज्ञान के बारे में यह आकर्षक जानकारी अधिकांशतः पांडित्यपूर्ण पत्रिकाओं और ग्रंथों में निहित है। इन सभी को औसत शिक्षित पाठक को सुलभ बनाने के लिए हमें अन्य विद्वानों की आवश्यकता है। यह छोटी सी पुस्तक स्पष्टता और लालित्य के साथ सटीक रूप से इसी उद्देश्य की पूर्ति करती है। इसके लेखक आलोक कुमार (alok.kumar@oswego.edu) एक प्रतिष्ठित प्रोफेसर हैं जिनके

नाम विद्वत्तापूर्ण प्रकाशनों की एक लंबी सूची है। तकनीकी विज्ञान के अलावा जिसमें वह अनुसंधान और पाठन कार्य करते हैं उन्होंने विज्ञान के इतिहास का भी गहन अध्ययन किया है और हिन्दू एवं अरब विज्ञान की विद्वत्तापूर्ण जांच पड़ताल की है। उनकी विज्ञान की दृष्टि व्यापक और सार्वभौमिक है।

पुस्तक की शुरुआत ज्ञान प्राप्ति की हिन्दू पद्धति से होती है, जो स्पष्ट रूप से दिखाती है कि इस संस्कृति में सत्य की खोज हमेशा श्रद्धा के साथ देखी जाती रही है और शिक्षक हमेशा सम्मान का पात्र रहा है। ज्ञान का प्रत्येक अन्वेषण परम वास्तविकता के बोध के व्यापक ढांचे और उस तलाश में निहित मानव चेतना की प्रासंगिकता को लेकर किया जाता रहा है।

पुस्तक में संख्या प्रणाली, ज्यामिति और त्रिकोणमिति के साथ-साथ पाई (π) के हिन्दू मूल्यांकन जैसे विषयों पर प्राचीन हिन्दू विचारकों के योगदानों की क्रमबद्ध चर्चा की गई है। अन्य स्थानों की तरह यहाँ भी पुस्तक बताती है कि कैसे हिंदू गणित फैला और उसने अन्य सभ्यताओं के ढांचे को प्रभावित किया। इस संदर्भ में जिस बात की ओर लेखक इंगित नहीं करता है, वह यह है, कि यद्यपि विभिन्न हिंदू (संस्कृत) ग्रंथों का चीनी, लैटिन, अरबी, फारसी और यूरोपीय भाषाओं में अनुवाद किया गया था लेकिन किसी भी अन्य संस्कृति के शास्त्रीय ग्रंथ का अनुवाद संस्कृत में नहीं किया गया। इससे इस बात की व्याख्या होती है कि प्राचीन हिंदू लेखन में विदेशी लेखकों के संदर्भ अथवा उनके योगदान की स्वीकार्यता क्यों नहीं है।

प्राचीन हिंदू विज्ञान में, कैलेंडर और ब्रह्माण्ड विज्ञान के संदर्भ में, विस्तार से हिंदू खगोल विज्ञान पर विचार किया गया है और, साथ ही उज्जैन की प्रसिद्ध वेधशाला (जिसे “प्राचीन विश्व का ग्रीनविच” कहा जाता है) का भी उल्लेख है। इसमें यह भी बताया गया है कि हिंदू खगोल विज्ञान मध्य पूर्व, चीन और यूरोप में कैसे फैला। अन्य विषय जिन पर विस्तार से चर्चा की गई है उनमें शामिल हैं: अंतरिक्ष, समय और पदार्थ (भौतिकी), खनन और धातु विज्ञान (रसायन विज्ञान) की धारणाएं तथा पारिस्थितिकी, चिकित्सा, और प्राचीन भारत में सर्जरी।

यह पुस्तक विज्ञान के लिए प्राचीन हिंदुओं के योगदान के पूरे परिसर को एक क्रमबद्ध, सुव्यवस्थित और पर्याप्त अच्छे तरीके से संदर्भ देते हुए प्रस्तुत करती है। इससे यह भी पता चलता है कि भारत में उभरे विचारों और अंतर्दृष्टि को दुनिया के कई अन्य हिस्सों में, विशेष रूप से यूरोप और अमेरिका में, विचारकों द्वारा कैसे सराहा और फैलाया गया था। प्रस्तुति की शैली और स्तर किसी भी शिक्षित व्यक्ति की पहुंच के भीतर है।

इस प्रकार यह पुस्तक इस विषय पर लगातार बढ़ते हुए साहित्य में एक अत्यंत मूल्यवान योगदान है। अपनी पुस्तक “प्राचीन काल से आधुनिक काल तक गणितीय विचार” (1990) में मॉरिस क्लाइन ने कहा है, “यह काफी हद तक निश्चित है कि हिंदुओं ने अपने स्वयं के योगदान के महत्व की सराहना नहीं की है।” उनकी यह बात सच है या नहीं पर इतना तो तय है कि क्षेत्र के विशेषज्ञों के अतिरिक्त, दुनिया के लोग प्राचीन हिंदुओं द्वारा विकसित वैज्ञानिक योगदान की पूरी तरह से सराहना नहीं करते हैं। यह पुस्तक इस गलतफहमी को दूर करने के लिए बहुत कुछ करती है। कोई आश्चर्यचकित हो सकता है कि लेखक ने इसका वर्णन भारतीय विज्ञान के बजाए हिंदू विज्ञान के रूप में क्यों किया : समीक्षक संभवतः इसे भारतीय विज्ञान कहना अधिक पसंद करता क्योंकि जैन और बौद्ध विचारकों का भी प्राचीन भारतीय विज्ञान के कुछ गहन विचारों में योगदान रहा था। लेकिन लेखक ने प्रस्तावना में स्पष्ट रूप से पुस्तक के शीर्षक के चयन का कारण बताया है। यद्यपि हिंदू शब्द स्वयं विदेशी मूल का है, किन्तु आधुनिक भारत में यह एक जातीय रूप से संवेदनशील विशेषण बन गया है। पिछले कई दशकों से, अब वहाँ आधुनिक भारत की हिंदू जड़ों पर

जोर देने के लिए आंदोलन किए जाते रहे हैं। आधुनिक पश्चिम ईसाई नहीं है, लेकिन इसकी सांस्कृतिक जड़ें काफी हद तक ईसाई हैं। ऐसा ही आधुनिक भारत के साथ है: एक आधुनिक धर्मनिरपेक्ष लोकतंत्र जिसकी जड़ें निश्चित रूप से हिंदू हैं।

उस दृष्टिकोण से देखें तो पुस्तक का शीर्षक उचित ही प्रतीत होता है।

पुस्तक निश्चय ही आधुनिक वैश्विक सभ्यता में एक प्रमुख कारक के रूप में पाठक की समझ और विज्ञान- दृष्टि को बढ़ाएगी।

— वी. वी. रमन, भौतिकी और मानविकी के प्रोफेसर (एमेरिटस)
Rochester Institute of Technology, Rochester NY, USA; vvrsp@rit.edu

(अनुवादक: राम शरण दास गुप्त, rsgupta_248@yahoo.co.in)

डॉ. वेद प्रकाश चौधरी द्वारा भारत में प्रकाशित, पुस्तक की समीक्षा

ANCIENT HINDU SCIENCE - Its Transmission and Impact on World Cultures

Author: Alok Kumar, Publisher: JAICO BOOKS - India; www.jaicoBooks.com

यह पुस्तक संक्षिप्त किन्तु व्यापक और भलीभांति शोधपूर्वक रचित कृति है, जो विज्ञान के मूलधार तत्त्वों का विहंगम इतिहास प्रस्तुत करती है।

प्रोफेसर डॉ. आलोक कुमार ने कई दशाब्दियों के शोध और समन्वय के उपरांत एक अति उत्कृष्ट कृति की रचना की है। उन्होंने संपूर्ण प्राचीन विश्व (भारत, अरब, यूरोप और चीन) के सभी ऐतिहासिक स्रोतों से बहुल मात्रा में लेख, प्रलेख, प्रमाण, दस्तावेजों और उपकरणों के चित्रों का संग्रह किया है और उन पर व्यापक शोध और समन्वय किया है। वे प्राचीनतम हिंदू ग्रंथों का विवरण देते हैं जैसे, ऋग्वेद और ब्राह्मण (2000 ईसा पूर्व या उससे भी पहले), उपनिषद, शुल्बसूत्र, वैशेषिक सूत्र, सुश्रुत संहिता, चरक संहिता, कौटिल्य का अर्थशास्त्र, आचार्य पिंगला का छंद सूत्र (पाँचवीं शताब्दी ईसा पूर्व), से लेकर प्रतिभाशाली आचार्य आर्यभट्ट के प्रसिद्ध ग्रंथ आर्यभट्टिया (पाँचवीं शताब्दी वर्तमानकालीन) तथा बक्षाळी पांडुलिपि (7 वीं शताब्दी वर्तमानकालीन)

और विभिन्न पुराणों (पहली शताब्दी से लेकर 11 वीं शताब्दी तक वर्तमानकालीन) सभी ग्रंथों से उन्होंने प्राचीनतम तथ्यों का संग्रह, शोध और समन्वय इस पुस्तक में प्रस्तुत किया है।

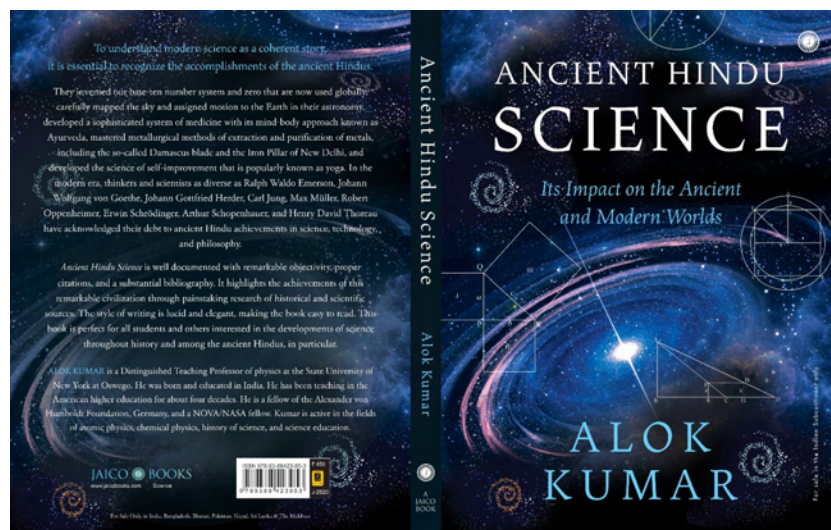
डॉ. आलोक कुमार ने कई वर्षों के अथक परिश्रम के उपरांत, विस्तार पूर्वक मूल आलेखों और उपकरणों को चित्रित करते हुए, विज्ञान और गणित की उन मूलभूत अवधारणाओं के इतिहास की सीमा को स्थापित किया है जिनको प्राचीन हिंदुओं ने विश्व में सर्वप्रथम, अन्य देशों से हजारों वर्ष पूर्व खोज लिया था। यह संक्षिप्त पुस्तक गणित, भौतिकी, रसायन विज्ञान, खनन और धातु विज्ञान, खगोल विज्ञान, चिकित्सा, जैविक विज्ञान, आयुर्वेद के साथ-साथ मन-बुद्धि-शरीर (ज्ञानेन्द्रियों) के एकीकरण हेतु ऋषियों द्वारा स्थापित तकनीकों (योग, प्राणायाम, ध्यान) के मूलभूत सृजन सूत्रों को प्रदर्शित करती है जिन्हें प्राचीन हिंदुओं ने 2000 ई.पू. से 500 ई.पू. या उस से भी पहले खोज लिया था।

अरब विद्वानों ने तो सरल भावना से हिन्दुओं से प्राप्त ज्ञान को हिन्दू विज्ञान की ही संज्ञा दी थी (उदहारण: अरब विद्वान तो भारत की दश अंकों की प्रणाली को हिन्दू नूमरल्स ही कहते थे, परन्तु पाश्चात्य देशों ने इनको अरब नूमरल्स की संज्ञा दी। भारतीय वैज्ञानिक जो पाश्चात्य विश्वविद्यालयों में काम करते

हैं उन्होंने जब इस तथ्य पर प्रकाश डाला (और जोर डाला) तब पिछली शताब्दी के अंतिम चरण में विश्व के सभी वैज्ञानिकों ने यह स्वीकार किया कि वास्तव में यह हिन्दू नूमरल्स हैं, और सारे विश्व में अब यह हिन्दू नूमरल्स कहलाते हैं ।

विज्ञान का इतिहास जिसे सारी दुनिया जानती है, और जिसे सारे विश्व के छात्रों को पढ़ाया जाता है, औपनिवेशिक (कॉलोनिअल) काल के दौरान पश्चिमी विद्वानों द्वारा लिखा गया था, जिन्हें हिन्दू ऋषियों और प्राचीन वैज्ञानिकों के संस्कृत भाषा में लिखे मूल स्रोतों को समझने की क्षमता नहीं थी और न उन्हें इनके प्रति कोई आदर के साथ सीखने की भावना ही थी। इसलिए प्राचीन हिंदुओं के योगदान को उन्होंने तिरस्कृत किया और अपने लेखों, पुस्तकों और अन्य रचनाओं से यह स्थापित किया कि विज्ञान की उत्पत्ति यूरोप में पिछले 500 वर्ष में हुई है, और सारे विश्व ने इसे पश्चिमी स्रोतों से ही सीखा है। आधुनिक वैज्ञानिक इतिहास की यही विडम्बना है।

एक वैज्ञानिक होने के नाते मैं जीवन भर यह सोचता रहा कि यदि हमारे ग्रंथों में ऐसे वैज्ञानिक तथ्यों का वर्णन है तो किसी भारतीय वैज्ञानिक ने इस पर शोध करके एक ऐसी प्रामाणिक पुस्तक क्यों नहीं लिखी जिसको पढ़ कर हमें यह ज्ञात हो सके कि क्या तथ्य है (और क्या मिथ्या कहानियाँ जुड़ गई हैं हमारे समाज के सोच समझ में क्योंकि लम्बे समय से कोई प्रामाणिक पुस्तक उपलब्ध नहीं थी।), अंततोगत्वा, अब एक ऐसी पुस्तक उपलब्ध है (Amazon-in से)।



डॉ. आलोक कुमार भौतिक विज्ञान के एक प्रतिष्ठित (डिस्टिंग्विश्ड) प्रोफेसर की उपाधि से पुरस्कृत किये गए हैं, न्यूयॉर्क राज्य की एक राजकीय यूनिवर्सिटी, में जो न्यूयॉर्क के उत्तरी भाग में ओस्वीगो नगर में स्थित है।

वह संपूर्ण विश्व के विज्ञान के इतिहास पर कई पुस्तकों के लेखक और सह-लेखक भी हैं, जैसे

- Science in the Medieval World, 1991 and 1996
- Sciences of the Ancient Hindus: Unlocking Nature in the Pursuit of Salvation, 2014
- A History of Science in World Cultures: Voices of Knowledge, 2015
- Ancient Hindu Science: Its Transmission and Impact on World Cultures, 2019.

समीक्षक —डॉ. वेद प्रकाश चौधरी (Ved.chaudhary@gmail.com), अध्यक्ष, ESHA www.ESHAusa.org

प्रतिक्रियाएं / Feedbacks

सर्वप्रथम हम आपको धन्यवाद देना चाहेंगे की आप अति उत्साह के साथ विज्ञान जैसे गम्भीर विषय को अत्यंत सरल रूप में विज्ञान प्रकाश जर्नल के माध्यम से पाठकों के बीच में प्रस्तुत कर रहे हैं। हिंदी माध्यम में, विज्ञान जैसे संकाय में, किसी चुने हुए विषय पर समस्त तकनीकी पहलुओं को सम्मिलित करते हुए पाण्डुलिपि तैयार करना एवं पाण्डुलिपि के समस्त शीर्षकों के बीच सामंजस्य बनाये रखना एक बड़ी चुनौती के रूप में रहा। इस दौरान हमने विज्ञान प्रकाश जर्नल में दिए गए दिशा निर्देशों का व्यापक अध्ययन किया। जर्नल में प्रकाशित कई शोध पत्रों को बारीकी से पढ़ा और फिर हमने अपने रिसर्च पेपर को हिंदी में लिखने का कार्य शुरू किया। विज्ञान प्रकाश जर्नल में प्रकाशित शोध पत्रों में गुणवत्ता से कोई भी समझौता नहीं किया गया है, समस्त प्रकाशित शोध पत्र उच्च श्रेणी के हैं। शोध पत्र की गंभीर रूप से समीक्षा की गयी। विज्ञान प्रकाश जर्नल में छपे लेखों में पिक्चर्स एवं डाटा, रंगीन होने के कारण प्रस्तुत किये गए बिंदुओं को और अधिक स्पष्ट कर देते हैं जो कि अन्य जर्नल से भिन्न है। यह इस जर्नल की एक विशेष पहचान है। विज्ञान प्रकाश जर्नल द्वारा पाण्डुलिपियों का स्वीकृत करने से पहले गहन परीक्षण एवं कड़ी समीक्षा प्रक्रिया गुजारना इस जर्नल की विशेषता है। हमें आशा है की हमारा शोध पत्र अति शीघ्र विज्ञान प्रकाश जर्नल के आगामी अंक में प्रकाशित होगा।

*Authors: Dr. Shilpy Shakya & Dr. Bindhya Chal Yadav; राजकीय महाविद्यालय फतेहाबाद, आगरा;
shilpy.shilpy@gmail.com , dradarshmangal1@hotmail.com*

हिंदी में तकनीकी अनुसंधान पत्र लेखन चुनौतीपूर्ण है, लेकिन यह रचनात्मक है। विज्ञान प्रकाश जर्नल के माध्यम से मुझे पहली बार हिंदी में तकनीकी शोध पत्र लिखने का अवसर मिला। इस (विज्ञान प्रकाश) जर्नल में रिसर्च पेपर के सबमिशन के बाद मुझे हिंदी में रिसर्च पेपर लिखने की बारीकियां पता चली। यहाँ रिसर्च पेपर क्वालिटी से कोई समझौता नहीं है, अतः गुणवत्ता की कसौटी पर अति उत्तम है। अपने शोध पत्र को इस जर्नल में प्रकाशित करने का मेरा अनुभव काफी अच्छा रहा।

*Author: Ms Sofia Goel & Dr. Sudhansh Sharma, SOCIS, IGNC, New Delhi;
sofiagoel@gmail.com , sudhansh@ignou.ac.in*

विज्ञान प्रकाश पत्रिका में लिखने और रिव्यू करने का अवसर प्राप्त हुआ। हिंदी भाषा में शोध पत्र लिखने का अनुभव इतना सरल और सहज रहा कि अपने स्कूल के दिनों की विज्ञान की कक्षाएं पुरस्मृति में आ गयी। विज्ञान के कुछ अंग्रेजी शब्दों के हिंदी प्रतिरूप ज्यादा लोकप्रिय नहीं होते जिससे कि भाषा में लचीलापन नहीं आ पाता, परन्तु उन्हें अपने लोकप्रिय रूप में ही ग्रहण कर लेने की सम्भावना के फलस्वरूप अच्छी तरह से अपने विचार व्यक्त किये जा सके। अन्य लेखकों के शोध पत्रों की रिव्यू करते हुए तीन दिशाओं में प्रयत्न करने होते हैं : १) शोध का डोमेन २) लेख की भावनाएं और ३) हिंदी भाषा में लेखन कौशल। ऐसा अहसास हुआ कि ज्यादातर लेखक हिंदी भाषा में लेखन के लिए मूलभूत लिखने के बजाय ओपन सोर्स ट्रांसलेशन सिस्टम का उपयोग करते होंगे जिससे कि भाषा में कुछ अटपटापन रह जाता है। खैर, हिंदी भाषा में शोध पत्र लिखने के लिए प्रेरित होना और प्रयास करना ही अपने आप में एक बड़ी उपलब्धि है।

हर एक प्रकाशन के साथ, विज्ञान प्रकाश के कॉलम में क्रमशः उत्तरोत्तर वृद्धि और विभिन्न प्रकार के लेखों का समावेश निश्चय ही रूप से प्रभावकारी है। मुझे इस मिशन से जुड़ने का सौभाग्य प्राप्त हुआ। सादर धन्यवाद !

Author & Reviewer: Dr. Priyanka Jain, C-DAC; priyankaj@cdac.in

UGC-CARE listed research journal VIGYAN PRAKASH
Vol. 18 , No. 3-4, July-December 2020 : LIST OF REVIEWERS

- **Dr. Ruchi Patel**
Assistant Professor, Medicaps University, Indore,
ruchipatel294@gmail.com
- **Dr. Richa Singh**
Assistant Professor, Senior Grade-I, VIT Chennai
richasingh.bv@vit.ac.in
- **Dr. Deepti Chopra**
Assistant Professor, Lal Bahadur Shastri Institute
of Management, Delhi
deeptichopra11@yahoo.co.in
- **Prof. G. S. Lehal**
Dept. of Computer Science, Punjab University
gslehal@gmail.com
- **Dr. Priyanka Jain**
Associate Director, C-DAC , New Delhi
priyankaj@cdac.in
- **Mr. Syed Mohammad Ahmad,**
Chief Technology Officer (CTO) & Consultant:
sayedmahmad@gmail.com
- **Dr. Rajesh Vasita**
Assistant Professor, School of Life Sciences,
Central University of Gujarat;
rajesh.vasita@gmail.com
- **Dr. Yogendra Kumar**
Associate Professor, Department of Botany, Govt
Degree College Nanauta, Saharanpur
Ykpanchal4ever@gmail.com
- **Dr. Supriya Tiwari**
Assistant Professor, Department of Botany,
Banaras Hindu University, Varanasi
supriyabhu@gmail.com
- **Prof. Anju Khandelwal**
Dept. of Mathematics, SRMS Institute of
Engineering & Technology, Bareilly
dranju07khandelwal@gmail.com
- **Prof. Gokul Chandra Sharma**
Retd. Prof. of Mathematics,
Dr. B R Ambedkar Univ. Agra;
gokulchandra5@gmail.com
- **Dr. Omkar Singh Lodhi**
Assistant Professor, Dept. of Computer Science,
Maharaja Agrasen College, University of Delhi
omkarsingh@mac.du.ac.in ; omkarsinghlodhi@gmail.com
- **Dr. Pramod Kumar Sagar**
Ph.D.(CSE), MTech, MCA, Asst. Prof. (Sr.G),
Dept. of Computer Applications, SRM University,
Delhi NCR Campus, Modinagar.
pramodks@srmist.edu.in ; pramodsagar.srm@gmail.com
- **Dr. Rakesh Prasad Badoni, Ph.D.**
Assistant Professor, Xavier School of Computer Sc.
and Engn., Xavier University, Bhubaneswar, Odisha
rakeshbadoni@gmail.com
- **Dr. Apoorva Mishra**
Asst. Professor, IIIT, Pune
apoorvamish1989@gmail.com
- **Dr. Vijay Shankar Tripathi**
Professor, MNNIT, Allahabad
vst@mnnit.ac.in
- **Dr. Mohd Gulman Siddiqui**
Asst. Prof., Banasthali Vidyapith
mohd gulman@gmail.com
- **Dr. Ashok Chandra, D.Sc.**
Former, Wireless Advisor to Govt. of India, Expert,
International Telecommunications Union (ITU),
drashokchandra@gmail.com
- **Dr. Priyanka Jain**
Associate Director, C-DAC, New Delhi
priyankaj@cdac.in
- **Dr. Sayed M Ahmad**
Chief Technology Officer & Consultant
Abul Fazal Part-1, New Delhi 110025
sayedmahmad@gmail.com
- **Prof. K. K. Mishra**
Homi Bhabha Center for Science Education
TIFR, Mumbai 400088
kkm@hbcse.tifr.res.in
- **Prof. Oum Prakash Sharma**
Director, NCIDE, IGNOU, Delhi
opsharma@ignou.ac.in
- **Prof. P. Jha**
Director, Cultural Informatics Lab, IGNCA
pjh@ignca.nic.in
- **Shri Ram Sharan Das**
4/49, Vaishali, Ghaziabad-201010, U.P.
rsgupta_248@yahoo.co.in
- **Prof. Om Vikas**
Hon. Advisor, Bharatiya Vidya Bhavan, Delhi
dr.omvikas@gmail.com

Vision: High-Tech research to reach out widely promoting inclusive innovation and entrepreneurship

Mission: Publication of quality research articles in Hindi in *Sciences* (Physics, Chemistry, Mathematics, Bio-science, Medical Science, AYUSH, Management Science, Agriculture and Environment), *Engineering* and *Technology*, and promoting creative ideas for innovation, incubation and entrepreneurship.

Submission: Title, Author affiliation, abstract and keywords be in both Hindi and English, and references as they are originally referred to. Overview Articles and research papers must be **original without plagiarism**. Authors need to mention three or more subject experts also from different institutions to review the submitted article. Articles may be submitted to Editor@VigyanPrakash.in

रवीन्द्रनाथ टैगोर रचित गीतांजलि से

चित्तं येथा भयशून्या उच्च येथा शिर,
ज्ञान येथा मुक्तं, येथा गृहेर प्राचीर
आपन प्राङ्गणतले दिवस-शर्बरी,
बसुधारे राखे नाइ खण्ड क्षुद्रकरि,

येथा वाक्य हृदयेर उतससृथ हत
उच्छसियया उठे, येथा निर्बारित स्रोते,
देशे देशे दिशे दिशे कर्मधारा धाय
अजस्र सहस्रविध चरितार्थताय,

येथा तुच्छ आचारेर मरु-बालु-बाशि
बिचारेर स्रोतःपथ फेले नाइ ग्रासि -
पौरुषेरे करेनि शतधा, नित्य येथा
तुमि सर्व कर्म-चिन्ता-आनन्देर नेता,

निज हस्ते निर्दय आघात करि पितः,
भारतेरे सेइ स्वर्गे करो जागरित //

Devanagari transliteration

चित्तं जेथा भयशून्य, उच्च जेथा शिर,
ज्ञान जेथा मुक्त जेथा गृहेर प्राचीर
आपन प्राङ्गणतले दिवस शर्बरी
बसुधारे राखे नाइ खण्ड क्षुद्रकरि,
जेथा वाक्य हृदयेर उत् समुखहते
उच्छसियया उठे, जेथा निर्बारित स्रोते
देशे देशे दिशे दिशे कर्मधारा धाय
अजस्र सहस्रविध चरितार्थताय
जेथा तुच्छ आचारेर मरुबालुराशि
बिचारेर स्रोतःपथे फेले नाइ ग्रासि,
पौरुषेरे करे नि शतधा, नित्य जेथा
तुमि सर्व कर्म चिन्ता आनन्देर नेता—
निज हस्ते निर्दय आघात करि पितः,
भारतेर सेइ स्वर्गे करो जागरित ।

भावानुवाद -

कवि शिवमंगल सिंह “सुमन” द्वारा

जहां चित्त भय से शून्य हो
जहां हम गर्व से माथा ऊंचा करके चल सकें
जहां ज्ञान मुक्त हो
जहां दिन रात विशाल वसुधा को खंडों में विभाजित कर
छोटे और छोटे आंगन न बनाए जाते हों
जहां हर वाक्य हृदय की गहराई से निकलता हो
जहां हर दिशा में कर्म के अजस्र नदी के स्रोत फूटते हों
और निरंतर अबाधित बहते हों
जहां विचारों की सरिता
तुच्छ आचारों की मरु भूमि में न खोती हो
जहां पुरुषार्थ सौ सौ टुकड़ों में बंटा हुआ न हो
जहां पर सभी कर्म, भावनाएं, आनंदानुभूतियाँ
तुम्हारे अनुगत हों
हे पिता, अपने हाथों से निर्दयता पूर्ण प्रहार कर
उसी स्वातंत्र्य स्वर्ग में इस सोते हुए भारत को जगाओ

Translation by
Rabindranath Tagore himself
(Geetanjali English Version)

Where the mind is without fear and
the head is held high;
Where knowledge is free;
Where the world has not been broken up
into fragments by narrow domestic walls;
Where words come out from the depth of truth;
Where tireless striving stretches
its arms towards perfection;
Where the clear stream of reason
has not lost its way into the
dreary desert sand of dead habit;
Where the mind is led forward by thee
into ever-widening thought and action—
Into that heaven of freedom, my Father,
let my country awake.

विज्ञान प्रकाश : विज्ञान एवं प्रौद्योगिकी रिसर्च जर्नल

VIGYAN PRAKASH : Research Journal of Science & Technology

www.VigyanPrakash.in