

अध्याय 6

कृत्रिम बुद्धि तथा

अभिकलनात्मक व्याकरण

6.1 कृत्रिम बुद्धि

6.2 व्याकरण तथा अभिकलनात्मक व्याकरण

6.3 अनुवादक तथा दुभाषिया

6.3.1 अनुवादक

6.3.2 दुभाषिया

6.3.3 हिंदी कंप्यूटर में अनुवाद की स्थिति

- (i) मंत्र (ii) अनुवादक 2.0 (iii) नामानुवाद (iv) लिप्यंतरण
- (v) शब्दानुवाद (vi) पर्यायवाची कोश (vii) राजभाषा हिंदी कोश
- (viii) शक्ति - त्रिभाषी अनुवादक
- (ix) प्रदर्श 7-10

6.4 आवाज पहचानक/वाक् अभिज्ञान

6.4.1 विकास गाथा

6.4.2 आवाज पहचानक : कैसे और क्या

6.4.3 विभिन्न पहचानक/वाक् अभिज्ञान

- (क) इंटरफेस वायावाइस (ख) नेट कंटेंटर
- (ग) स्पीच रेकॉग्निशन (घ) ड्रेगन नेचुरली स्पीकिंग
- (ङ) वायावाइस - 10.0 (च) ऑफिस एक्स पी स्टैंडर्ड

6.4.4 विभिन्न आवाज पहचानक उत्पादों का विश्लेषण

6.4.5 प्रशिक्षण का महत्त्व

6.4.6 भविष्य

6.5 सोनोग्राफ

6.6 वृक्ष संलग्न व्याकरण (टैग)

6.6.1 संदर्भयुक्त व्याकरण (CAG) और टैग (TAG) में अंतर

6.6.2 टैग का रूपवाद

6.6.3 टैग में निहित एकीकरण और लक्षण-संरचना

6.6.4 टैग की स्थानिकता का व्यवहार-क्षेत्र

6.6.5 टैग के अंतर्गत संलग्नक प्रतिबंध

6.6.6 टैग और हिंदी व्याकरण, बीज वाक्यों का कोशीय अंतरण, अंतरण शब्दवृत्त

6.7 अध्याय 6 की संदर्भ सूची

कृत्रिम बुद्धि तथा अभिकलनात्मक व्याकरण

6.1 कृत्रिम बुद्धि

कृत्रिम बुद्धि एक ऐसा विज्ञान है, जिसके अंतर्गत मशीन को तर्क करने, समस्याओं का समाधान खोजने, स्वयं निर्णय लेने और परीक्षण व अशुद्धि प्रणाली (Trial and Error Method) से स्वयं सीख लेने योग्य बनाया जाता है। विशेषज्ञ प्रणाली (Expert System), प्राकृतिक भाषा संसाधन (Natural Language Processing) आदि क्षेत्र कृत्रिम बुद्धि के अंतर्गत आते हैं। इनमें प्राकृतिक भाषा संसाधन का क्षेत्र सर्वाधिक महत्वपूर्ण है, क्योंकि इसकी सहायता से मशीन को प्राकृतिक भाषा समझने में सहायता मिलती है और इस प्रकार मानव मशीन के बीच बेहतर संवाद स्थापित हो जाता है। वस्तुतः कृत्रिम बुद्धि के वैज्ञानिकों ने भाषा की समस्या को संग्रहण की समस्या माना है, यही कारण है कि जहाँ एक ओर सामान्य भाषा वैज्ञानिक भाषा के पक्ष विशेष पर प्रकाश डालता है, वहाँ कृत्रिम बुद्धि के अंतर्गत वाक्यपरक, अर्थपरक और संदर्भपरक तत्वों को समन्वित रूप में देखने का प्रयास किया जाता है।

यदि ज्ञान में सामान्यीकरण की यह विशेषता न हो तो उसे निरूपित करने के लिए अधिक स्मृति-कोश की जरूरत होगी, इसे प्राप्त करने में समय भी अधिक लगेगा और इसे अद्यतन भी बनाए रखना होगा।

- इसे लोग जिस रूप में समझते हों, उसी रूप में प्राप्त किया जाना चाहिए, हालांकि कई क्रमादेशों में अधिकांश द्वारा स्वचालित रूप में ही अनेक साधनों से एकत्र कर लिया जाता है।
- अशुद्धियों को और विश्व के अपने परिदृश्य को ठीक करने के लिए इसे सरलता से संशोधित किया जा सकता हो।
- पूर्णतः शुद्ध या पूर्ण न होने पर भी इसका उपयोग अधिकांश स्थितियों में किया जा सकता हो।

कृत्रिम बुद्धि के अंतर्गत प्राकृतिक भाषा संसाधन

प्राकृतिक भाषा के पाठ को समझना इतना सरल नहीं है। इसके अनेक कारण हैं। सर्वप्रथम इसके लिए विपुल मात्रा में वास्तविक विश्वज्ञान के निरूपण और प्रकलन (Manipulation) की आवश्यकता है। दूसरे, उस भाषा के वाक्य-विन्यास और शब्दावली को समझने की जरूरत है।

वस्तुतः प्रत्येक ज्ञान संरचना (Knowledge Structure) एक ऐसी डाटा संरचना (Data Structure) है, जिसके अंतर्गत किसी समस्या विशेष के व्यवहार-क्षेत्र (Domain) से संबंधित ज्ञान को संगृहीत किया जा सके। ये संरचनाएं व्यवहार-क्षेत्र के भीतर की चीजों को ही वस्तुओं (Objects) या घटनाओं (Events) के रूप में दर्शाती हैं। इनका महत्व इसलिए भी है कि ये एक ऐसा उपाय सुझाती हैं, जिनकी सहायता से चीजों के समान रूप

से उभरने वाले स्वरूप से संबंधित सूचनाओं को निरूपित किया जा सके। इस प्रकार के वर्णन को विवरणिका (Schema) कहा जा सकता है।

विवरणिका अतीत की प्रतिक्रियाओं या अतीत के अनुभवों का एक ऐसा सक्रिय संगठन है, जो अनुकूलित (Adapted) जैविक अनुक्रिया (Organised Response) के रूप में परिचालित हुआ माना जाना चाहिए (Barlett, 1932, पृ.201)।

विवरणिकाओं का प्रयोग करते हुए हम इस तथ्य का समुपयोजन (Exploit) कर सकते हैं कि वास्तविक विश्व यादृच्छिक (Random) नहीं है। अनेक प्रकार की विवरणिकाएँ कृत्रिम बुद्धि के क्रमादेशों के अंतर्गत उपयोगी सिद्ध हुई हैं, जैसे :

- फ्रेम (Frames) जिनका उपयोग किसी वस्तुविशेष, जैसे मेज आदि की विशेषताओं के संकलन को दर्शाने के लिए किया जाता है।
- स्क्रिप्ट्स जिनका उपयोग घटनाओं के सामान्य क्रम को दर्शाने के लिए किया जाता है, जैसे कि कोई ग्राहक जब रेस्तरां में जाता है तो क्या-क्या होता है।

इन दोनों ज्ञान-संरचनाओं का कृत्रिम बुद्धि के अनेक क्रमादेशों (Programs) में व्यापक रूप से उपयोग किया जाता है।¹

6.2 व्याकरण तथा अभिकलनात्मक व्याकरण

कृत्रिम बुद्धि के विकास के कारण आज मशीनी अनुवाद के क्षेत्र में अनेक अनपेक्षित सफलताओं के द्वार खुलने लगे हैं, लेकिन पूर्णतः स्वचालित मशीनी अनुवाद प्रणाली का विकास अभी तक विश्व में कहीं भी संभव नहीं हो पाया है। जिन प्रणालियों को पूर्णतः स्वचालित माना भी जाता है उनकी भी अपनी सीमाएँ हैं। उदाहरण के लिए कनाडा में विकसित TAUM METEO प्रणाली की शब्दावली, पदबंध, वाक्य संरचनाएँ और व्यवहार क्षेत्र केवल मौसम विज्ञान के बंधे-बंधाएँ विषय तक ही सीमित हैं। इसी प्रकार जापान में विकसित MU प्रणाली भी हवाई जहाज की नियमावली के सीमित दायरे में ही बंधी है।

भारतीय संदर्भ में आई.आई.टी. कानपुर द्वारा विकसित 'अनुसारक' निश्चय ही मशीनी अनुवाद की दिशा में मील का पत्थर साबित हो सकता है, लेकिन उसका दायरा भारतीय भाषाओं के बीच सीमित है। भारतीय भाषाओं में अंतर्निहित समानताओं का क्षेत्र बहुत व्यापक है। शब्दावली, वाक्य संरचना और सांस्कृतिक परिवेश की समानता के कारण इनमें परस्पर अनुवाद अपेक्षाकृत सरल ही माना जाता है।

भारतीय भाषाओं के बीच मुख्य बाधक तत्व है लिपि। ऊपर से देखने पर दसों भारतीय लिपियों में अनेक असमानताएँ दिखाई पड़ती हैं, किंतु ये सभी लिपियाँ ब्राह्मी लिपि से विकसित हुई हैं, जिसके फलस्वरूप उनमें अंतर्निहित समानताएँ भी हैं। सभी लिपियाँ ध्वन्यात्मक (Phonetic) हैं और सबका वर्णक्रम भी एक ही है। अभिकलनात्मक दृष्टि से (Computationally) इस समानता के कारण ही कंप्यूटर वैज्ञानिकों ने ब्राह्मी लिपि पर आधारित दसों भारतीय लिपियों के लिए समान कोड निर्धारित किया है। इसे ISCII 8 कहा जाता है। इलेक्ट्रॉनिकी विभाग की संस्तुति के आधार पर भारतीय मानक ब्यूरो ने इसे

मानकीकृत भी कर दिया है। समान इसकी कोड के कारण ही आज सभी भारतीय भाषाओं के लिए समान कुंजीपटल का विकास भी हो गया है। इसके फलस्वरूप भारतीय भाषाओं में परस्पर लिप्यंतरण स्वचालित रूप से संभव हो सका है। इस परिप्रेक्ष्य में अनुसारक प्रणाली का महत्त्व और भी बढ़ जाता है।

इस संदर्भ में यह बात भी काफी दिलचस्प है कि भारतीय भाषाओं में वाक्य संरचना के स्तर पर कुछ असमानताएं भी हैं, जिन्हें हिंदी के संदर्भ में क्षेत्रीय विविधताओं के रूप में स्वीकार किया जा सकता है और उनके कारण अर्थ का अनर्थ होने की आशंका नहीं होती। इस प्रकार की क्षेत्रीय विविधताओं के स्वीकार्य होने के कारण छोटी-मोटी अशुद्धियों की अनदेखी की जा सकती है। उदाहरण के लिए मराठी भाषा में हिंदी में अनूदित वाक्य (2) भी स्वीकार्य है।

वाक्य (1) महात्मा गांधी यांचा वाढ दिवस दोन अक्टूबरला आहे।

वाक्य (2) महात्मा गांधी इनका जन्म दिवस दो अक्टूबर को है।

हिंदी की प्रकृति के अनुसार इसका सही अनुवाद वाक्य (3) के रूप में होगा,

वाक्य (3) महात्मा गांधी का जन्म दिवस दो अक्टूबर को है।

यद्यपि अर्थ की दृष्टि से वाक्य (2) में कोई असंगति नहीं है, लेकिन वाक्य संरचना की दृष्टि से इसे पूर्णतः शुद्ध नहीं माना जा सकता, किंतु अर्थ प्रकट करने में यह वाक्य पूर्णतः सफल और सार्थक कहा जाएगा। इसमें संदेह नहीं कि इस प्रकार की असंगतियों को मानव अनुवादक द्वारा सरलता से दूर किया जा सकता है। थोड़े बहुत मानवीय हस्तक्षेप की व्यवस्था विश्व की अधिकांश प्रणालियों में की गई है, ताकि वाक्येतर तत्वों या प्रोक्ति के स्तर पर होने वाली अशुद्धियों का निराकरण मानवीय हस्तक्षेप द्वारा किया जा सके।

किंतु मशीनी अनुवाद की जटिलता उन भाषाओं के संदर्भ में काफी बढ़ जाती है, जब उनमें वाक्य-विन्यास के स्तर पर भारी विभिन्नता दिखाई पड़ती हो। उदाहरण के लिए अँग्रेजी-हिंदी को लिया जा सकता है। आज भारतीय संविधान में राजभाषा के रूप में हिंदी को प्रमुख दर्जा मिलने के बावजूद अँग्रेजी का वर्चस्व बना हुआ है। भारत सरकार द्वारा केंद्र सरकार के सभी विभागों में हिंदी अनुभाषा गठित किए गए हैं। हजारों की संख्या में हिंदी अनुवादकों की भर्ती की गई है, किंतु प्रत्येक प्रारूप मूलतः अँग्रेजी में तैयार होने के कारण मात्र सांविधिक दायित्वों के अनुपालन के लिए ही मूलतः अँग्रेजी में लिखित और मुद्रित लाखों पुस्तकों का हिंदी में साथ-साथ अनुवाद करना आवश्यक है। आज हजारों अनुवादकों की भर्ती के बावजूद मूल प्रारूपण के कारण अँग्रेजी पाठ पहले तैयार हो जाता है और तात्कालिक आवश्यकता के कारण उसे पहले जारी भी कर दिया जाता है और अंत में यह लाइन जोड़ दी जाती है 'हिंदी पाठ बाद में जारी कर दिया जाएगा' (Hindi version will follow)।

स्वचालित अनुवाद प्रणाली विकसित होने के बाद हिंदी पाठ भी अँग्रेजी के साथ-साथ वाक्य प्रति वाक्य तैयार हो जाएगा और मानव अनुवादक उसे देखकर जाँच कर लेंगे कि उसमें किसी प्रकार की अशुद्धि तो नहीं रह गई। इससे लाखों मानव घंटों की बचत होगी और अनुवाद भी अपेक्षाकृत अधिक शुद्ध होगा। इसके अलावा अँग्रेजी आज अंतरराष्ट्रीय भाषा होने

के साथ-साथ विज्ञान, प्रौद्योगिकी और व्यापार की भाषा भी है। हिंदी में अधुनातन ज्ञान-विज्ञान को यथाशीघ्र सुलभ कराने के लिए स्वचालित मशीनी अनुवाद की कोई प्रणाली अभी तक विकसित नहीं हो पाई है। इसका एक कारण कदाचित् यह है कि अंग्रेजी भाषा के सामान हिंदी भाषा का भाषिक विश्लेषण इस स्तर का नहीं हो सका है कि उसे मशीनी अनुवाद के लिए आधार बनाया जा सके।

हिंदी का वाक्यात्मक व्याकरण : 'हिंदी का वाक्यात्मक व्याकरण' में प्रो.सूरजभान सिंह (1985) ने लक्षण विश्लेषण के आधार पर यह सिद्ध करने का प्रयत्न किया है कि प्रत्येक शब्द के अंतर्गत उस भाषा के मूल संरचनात्मक लक्षण अंतर्निहित होते हैं। वस्तुतः यही लक्षण पदबंधों की वे भेदक विशेषताएं हैं जो उन्हें वाक्य के बाह्य और आंतरिक संरचना के स्तर पर अन्य सजातीय पदबंधों से अलग करती हैं। प्रो.सिंह की मान्यता है कि सामान्यतः वाक्य के प्रत्येक घटक को लक्षणबद्ध किया जा सकता है और उन्होंने कंप्यूटर की द्वयंक (Binary) पद्धति के अनुरूप उनमें विद्यमान विशिष्ट लक्षणों को ' + ' चिह्न द्वारा और उनके अभाव को ' - ' चिह्न द्वारा दर्शाया है। जैसे < + चेतन >, < - चेतन > आदि। प्रत्येक वाक्य सांचा लक्षणों का संविन्यास (Configuration) होता है।

ये लक्षण दो प्रकार के होते हैं, प्रकट और अप्रकट लक्षण। वचन, लिंग, काल, पक्ष आदि चिह्नों के माध्यम से प्रकट व्याकरणिक लक्षण सामान्य लक्षण कहलाते हैं। चेतन, मानव, मूर्त, गणनीय आदि जैसे लक्षण, जिनका संबंध बाह्य संरचना से न हो कर अर्थ तत्त्व से होता है अप्रकट लक्षण कहलाते हैं।

प्रो. सिंह ने इन्हीं लक्षणों के आधार पर हिंदी वाक्य सांचों को संरचना की दृष्टि से मोटे रूप में तीन प्रमुख वर्गों में रखा है : कोप्यूलान्ता वाक्य सांचे, को-वाक्य सांचे तथा क्रिया प्रधान वाक्य सांचे। इनके अंतर्गत हिंदी के 14 वाक्य सांचे बनते हैं। आर्थी, संरचनात्मक तथा अंतर्निहित लक्षणों के आधार पर इन बीज वाक्य सांचों के 45 उप सांचे बनते हैं। इस प्रकार हिंदी के सभी सरल वाक्य इन वाक्य सांचों में समाहित हो जाते हैं।

सामान्यतः वाक्य के प्रकारों का निर्धारण संरचना और अर्थ दो स्तरों पर किया जा सकता है। संरचना की दृष्टि से वाक्य के दो स्थूल भेद संभव हैं : सरल और असरल वाक्य। ऊपर बताए गए 14 प्रकार के वाक्य सांचे और उनके 45 उप वाक्य सांचे सरल वाक्यों की श्रेणी में आते हैं। अब यदि हम संरचना की दृष्टि से असरल वाक्यों का विश्लेषण करें तो उनके दो उपभेद किए जा सकते हैं, संयुक्त तथा मिश्र वाक्य।

संयुक्त वाक्य में आंतरिक संरचना के स्तर पर दो या अधिक स्वतंत्र उप वाक्य होते हैं, जिनमें परस्पर समानाधिकरण संबंध होता है अर्थात् वे अर्थ की दृष्टि से परस्पर आश्रित नहीं होते, इसके विपरीत मिश्र वाक्यों में यह संबंध आश्रय-आश्रित का होता है। उदाहरण के लिए वाक्य (1) और (2) देखें :

(1) रात सुनसान थी और चारों ओर अंधेरा था। (संयुक्त वाक्य)

(2) बाहर एक लड़का खड़ा है जो आपको ढूँढ़ रहा है। (मिश्र वाक्य)

इसी प्रकार अर्थ की दृष्टि से भी वाक्य के तीन प्रमुख वर्ग संभव हैं : कथनात्मक, आज्ञार्थक और मनोभावात्मक। देखें वाक्य (3), (4) और (5)

(3) राम खाना खा रहा है। (कथनात्मक वाक्य)

(4) राम को खाना खिलाओ। (आज्ञार्थक वाक्य)

मनोभावात्मक वाक्यों में विभिन्न वृत्ति रूपों का प्रयोग किया जाता है। संभावनार्थक, इच्छार्थक तथा विसमयादि बोधक वाक्य भी इसी के अंतर्गत आते हैं। देखें वाक्य (5), (6) और (7) :

(5) शायद अब वह न लौटे। (संभावनार्थक)

(6) आपकी यात्रा शुभ हो। (इच्छार्थक)

(7) अहा ! कितना सुंदर दृश्य है। (विस्मयादि बोधक)

वस्तुतः ये सभी वाक्य बीज वाक्यों से ही निर्मित होते हैं। परंपरागत व्याकरण और रूपांतरण व्याकरण दोनों यह स्वीकार करते हैं कि संयुक्त वाक्यों के दो उप-वाक्य होते हैं। रूपांतरण व्याकरण (Transformational Grammar) में इन्हें **आभ्यंतर वाक्य** या **संरचक वाक्य** (Constituent Sentence) कहते हैं। ये दोनों उप-वाक्य समानाधिकरण संबंध से जुड़े होते हैं लेकिन दो संरचक वाक्य आंतरिक संरचना से बाह्य संरचना तक आते-आते किस प्रकार संकुचित तथा संयोजित होकर एक वाक्य बन जाते हैं, इस रचना प्रक्रिया का नियम **रूपांतरण व्याकरण** में किया गया है। इसके अनुसार हिंदी के संयुक्त वाक्यों की रचना प्रक्रिया में अधिक से अधिक तीन रूपांतरण प्रक्रियाएं शामिल होती हैं :

(क) संयोजन प्रक्रिया :

वाक्य (1) रात सुनसान थी।

(2) चारों ओर अंधेरा था।

(3) रात सुनसान थी और चारों ओर अंधेरा था।

(ख) लोप प्रक्रिया :

यदि दो या अधिक संरचक वाक्यों में ऐसे समरूप अंश हों जो समान कोटि के हों और समान प्रकार्य करते हों तो दूसरे संरचक वाक्य के समरूप अंश का लोप हो जाता है।

वाक्य (1) मां सो रही है और बच्ची भी।

(2) मां सो रही है।

(3) बच्ची भी सो रही है।

(ग) रूपांतरण प्रक्रिया :

दूसरे संरचक वाक्य के असमान अंशों का पहले वाक्य में रूपांतरण हो जाता है। जहां वे पहले वाक्य के समकक्ष अंश के साथ संयोजन प्रक्रिया द्वारा संयोजित होते हैं।

वाक्य (1) लोहा धातु है।

(2) सोना धातु है।

(3) लोहा और सोना धातु है।

कुछ संयुक्त वाक्यों में एक प्रक्रिया काम करती है, कुछ में दो, कुछ में तीनों। मिश्र वाक्य में कम से कम दो उप वाक्य होते हैं। मुख्य/स्वतंत्र उप वाक्य और गौण/आश्रित उप वाक्य। मुख्य उप वाक्य वाक्य का मुख्य कथन होता है और आश्रित उप वाक्य कुछ समुच्चयबोधक अवयवों (जो, कि, ताकि आदि) द्वारा मुख्य उपवाक्य में जुड़े होते हैं।

वाक्य (1) बाहर एक लड़का खड़ा है।

(2) वह आपको दूँद रहा है।

(3) बाहर एक लड़का खड़ा है जो आपको दूँद रहा है।

मुख्य उपवाक्य बिना आश्रित उपवाक्य के भी प्रयुक्त होने की क्षमता रखता है, लेकिन आश्रित उपवाक्य अपने अर्थ की पूर्णता के लिए मुख्य उपवाक्य की आकांक्षा करता है। प्रकार्य की दृष्टि से आश्रित उपवाक्य तीन प्रकार के होते हैं।

(क) संज्ञा उपवाक्य :

उसने कहा कि गाड़ी छूट गई।

(ख) विशेषण उपवाक्य :

बाहर एक आदमी खड़ा है जो आपको दूँद रहा है।

(ग) क्रियाविशेषण उपवाक्य :

जब मैं कलकत्ता में था तो खुद खाना बनाता था।

अब प्रो. आर.एम.के.सिन्हा के निर्देशन में अँग्रेजी से हिंदी में मशीनी अनुवाद का प्रयास आई.आई.टी. कानपुर में शुरू कर दिया गया है। डॉ. हेमंत दरबारी (1991) ने डॉ. अरविंद जोशी द्वारा विकसित TAG (Tree Adjoining Grammar) रूपावाद (Formalism) के आधार पर संस्कृत भाषा के पद निरूपण के लिए व्याकर्ता का नाम के एक पदनिरूपित्र (Parser) का विकास किया था। इस पार्सर के आधार पर ही प्रो. सूरजभान सिंह और डॉ. दरबारी के सहयोग से विजय मल्होत्रा ने हिंदी-अँग्रेजी के कोशीय अंतरण (Lexical Transfer) की एक योजना तैयार की। डॉ. दरबारी का यह दावा था कि यदि हम हिंदी के वाक्यों का पद निरूपण (Parsing) करने में सफल हो जाते हैं, तो हिंदी-अँग्रेजी और अँग्रेजी-हिंदी के कोशीय अंतरण का मार्ग प्रशस्त हो जाएगा। तदनुसार डॉ. दरबारी के सहयोग से कंप्यूटर सोसाइटी ऑफ इंडिया द्वारा सन् 1993 में नई दिल्ली में आयोजित 'अक्षर' सेमिनार में प्रस्तुत प्रबंध के लेखक ने 'An efficient Parsing Technique for Hindi Language' नाम से एक आलेख प्रस्तुत किया।

हिंदी वाक्यों के सफल पद निरूपण के बाद डॉ. दरबारी ने English to Hindi Machine Translation through TAG नाम से एक आलेख प्रस्तुत किया, जिसमें व्याकर्ता के आधार पर पद निरूपित हिंदी वाक्यों के समकक्ष अँग्रेजी वाक्यों का पद निरूपण भी TAG के माध्यम से ही करने का प्रयास किया गया और स्रोत और लक्ष्य भाषा के रूप में अँग्रेजी और हिंदी शब्दों तथा पदबंधों के मूल वृक्ष (Elementary Tree) अंतरण शब्दवृत्त (Transfer Lexicon) में संकलित करके यह दर्शाने का प्रयास किया कि संरचना के स्तर

पर पर्याप्त भिन्नता के बावजूद इन दोनों भाषाओं का पदनिरूपण TAG के माध्यम से सरलता से किया जा सकता है।

चॉम्स्की (1957) ने प्रजनक व्याकरण के माध्यम से सर्वभाषा व्याकरण की एक संकल्पना भाषाविज्ञान के सम्मुख रखकर प्राकृतिक भाषा संसाधन का मार्ग खोल तो दिया लेकिन भाषाविशिष्ट प्रवृत्तियों के सांगोपांग विश्लेषण के अभाव में उन्हें निरंतर अपने सिद्धांतों में परिवर्तन, परिवर्धन और परिष्कार की आवश्यकता महसूस हुई।

मशीनी अनुवाद या कंप्यूटर साधित अनुवाद के स्थान पर 'कोशीय-अंतरण' शब्द का प्रयोग :

एक भाषा से दूसरी भाषा में अनुवाद की प्रक्रिया अत्यंत जटिल है। मशीनी अनुवाद के आधुनिक आरंभिक युग में कंप्यूटर विशेषज्ञ यह समझते थे कि अनुवाद संभवतः शब्द प्रति शब्द प्रतिस्थापना का ही कार्य है, किंतु बहुत जल्द उन्हें यह बोध हो गया कि अनुवाद कार्य संपन्न करने के लिए शब्द, अर्थ, वाक्य और संदर्भ सभी स्तरों पर विश्लेषण की आवश्यकता है। विजय मल्होत्रा द्वारा शब्द, अर्थ और वाक्य के स्तर पर मात्र उन्हीं कोशीय तत्वों के अंतरण का दावा किया गया है, जो शब्द विशेष की वृक्ष-संरचना में अंतर्निहित होते हैं। प्रत्येक भाषा की प्रवृत्ति भिन्न होने के कारण स्रोत और लक्ष्य दोनों भाषाओं के शब्दों के मूल संविन्यास में काफी भिन्नता हो सकती है। उदाहरण के लिए निम्नलिखित वाक्य देखें :

वाक्य (1) राम श्याम से मिलता है।

(2) Ram meets Shyam.

यदि हमने दोनों भाषाओं में प्रयुक्त 'मिलना' और 'meets' में अंतर्निहित लक्षणों को स्पष्ट नहीं किया तो निम्नलिखित वाक्य भी प्रजनित हो सकता है।

वाक्य (1) Ram meets Shyam.

संदर्भ या प्रोक्ति (discourse) के स्तर पर अनुवाद का प्रयास विजय मल्होत्रा ने अभी नहीं किया है, क्योंकि इसके लिए जिस स्तर पर विश्व ज्ञान (World knowledge) को संकलित करने की आवश्यकता है, वह अपने आप में एक दुष्कर कार्य है। उदाहरण के लिए निम्नलिखित वाक्य का कोशीय अंतरण इस पद्धति से फिलहाल संभव नहीं होगा।

वाक्य (1) उसने उसे मारा।

अँग्रेजी में इसके निम्नलिखित अनुवाद किए जा सकते हैं।

वाक्य (1) He beat him.

(2) She beat her.

(3) He beat her.

(4) She beat him.

वाक्येतर तत्वों के संधान के लिए टैग में कोई व्यवस्था नहीं है। हिंदी वाक्यों में भी इसका दायरा कुछ बीज वाक्यों और उनके उपवाक्य सांचों तक ही सीमित रखा गया है। संयुक्त वाक्य, मिश्र वाक्य, आज्ञार्थक वाक्य तथा मनोभावपरक वाक्यों के कोशीय अंतरण का इसमें प्रयास नहीं किया गया है।²

6.3 अनुवादक तथा दुभाषिया

(TRANSLATOR AND INTERPRETOR)

6.3.1 अनुवादक (Translator)

किसी एक भाषा में लिखित प्रोग्राम को दूसरी भाषा में समकक्ष प्रोग्राम में बदलने वाले प्रोग्राम को कंप्यूटर के क्षेत्र में अनुवादक या ट्रांसलेटर कहा जाता है। कंपाइलर तथा एसेंबलर इसके उदाहरण हैं। कंपाइलर उच्च-स्तरीय भाषा जैसे पास्कल में लिखे प्रोग्राम को मशीन कोड में अनूदित करता है और एसेंबलर एसेंबली भाषा में लिखे प्रोग्राम को मशीन कोड में अनूदित करता है।

प्रोग्रामिंग में किसी प्रोग्राम को एक भाषा से दूसरे में बदलने को अनुवाद कहा जाता है, जबकि कंप्यूटर ग्राफिक्स में डिस्प्ले पर प्रदर्शित 'स्थान' में किसी आकृति को, उल्टा घुमाए बिना, खिसकाने को अनुवाद कहते हैं।

अनूदित फाइल आँकड़ों वाली ऐसी फाइल को कहते हैं जो बाइनरी (8 बिट) प्रारूप से आस्की (7 बिट) प्रारूप में परिवर्तित की जा चुकी है। बिनहेक्स तथा ऑयन कोड दोनों बाइनरी फाइलों को आस्की में अनूदित करते हैं। जो सिस्टम (जैसे इ-मेल) प्रत्येक बाइट के आठवें बिट को परिरक्षित नहीं कर पाते हैं, उनमें से आँकड़ों को संप्रेषित करते समय ऐसे अनुवाद की आवश्यकता पड़ती है। अनूदित फाइल का प्रयोग करने से पहले उसे उसके बाइनरी प्रारूप में डिकोड किया जाना आवश्यक है।

अनुवाद-सेतु (Translating Bridge) दो भिन्न प्रकार के लेन-प्रोटोकॉल जैसे ईथरनेट और टोकन रिंग, आपस में अंतःसंबद्ध करता है। यह सामान्यतः बहुत जटिल उपकरण होता है। लेकिन सोर्स रूटिंग ट्रांसपरेट (एसआरटी) सेतु, ईथरनेट तथा टोकन रिंग दोनों सेतु पद्धतियों को जोड़ कर समस्या का समाधान करता है।

ट्रांस-लिस्प प्लस कंप्यूटर के लिए लिस्प का एक संस्करण है। इसे सोल्युशन सिस्टम, इंक., वेलेस्ले, एमए ने विकसित किया है। इसमें माइक्रोसॉफ्ट सी रूटीन को लिस्प लाइब्रेरी में एक फंक्शन के रूप में जोड़ने की सुविधा दी गई है।³

6.3.2 दुभाषिया (Interpreter)

'दुभाषिया' उच्च-स्तरीय भाषा में लिखे प्रोग्राम को कंप्यूटर के लिए स्वीकार्य तथा कार्यान्वयन के योग्य प्रारूप में बदलने का माध्यम है। यह ऐसा प्रोग्राम है जो उच्च-स्तरीय भाषा के प्रोग्राम - स्रोत प्रोग्राम - के कोड की प्रत्येक पंक्ति का विश्लेषण करता है तथा इसके बाद विनिर्दिष्ट कार्रवाई करता है। कंपाइलर के विपरीत, दुभाषिया कार्यान्वयन से पहले पूरे प्रोग्राम का अनुवाद नहीं करता है तथा कंपाइलर की तुलना में प्रोग्राम चलाने की इसकी गति अधिक धीमी है। लेकिन दुभाषिया लोड करने, स्रोत प्रोग्राम चलाने तथा स्रोत प्रोग्राम बदलने की प्रक्रिया का सरलीकरण करता है।

उच्च-स्तरीय भाषा का अनुवाद कुछ कंप्यूटरों पर दुभाषिए द्वारा किया जाता है तो अन्यो पर कंपाइलरों द्वारा किया जाता है। उदाहरणार्थ, बेसिक का अनुवाद दुभाषिए द्वारा किया जाता है। अधिक जटिल भाषाओं को कंपाइल किया जाना आवश्यक होता है।

जिस प्रोग्रामिंग भाषा को कार्यान्वित करने के लिए दुभाषिए की आवश्यकता पड़ती है, उसे **दुभाषियात्मक भाषा (Interpretive Language)** कहा जाता है।

जिस भाषा में प्रोग्राम को अनुदित करके कार्यान्वयन के योग्य बनाया जाता है तथा कार्यान्वयन से पहले पूरा का पूरा अनुदित (कंपाइल) किए जाने के स्थान पर एक बार में एक वाक्य का कार्यान्वयन किया जाता है, उसे **दुभाषित भाषा (Interpretated Language)** कहा जाता है। बेसिक, लिस्प, तथा एपीएल सामान्यतः दुभाषित भाषाएँ हैं। वैसे बेसिक को कंपाइल किया जा सकता है।⁴

6.3.3 हिंदी कंप्यूटर में अनुवाद की स्थिति

(i) मंत्र (MANTRA - Machine Assisted Translation)

मंत्र मशीन की सहायता से अनुवाद करने वाला यंत्र है, जो अंग्रेजी पाठ को हिंदी में निर्दिष्ट दायरे में अनुदित करता है। मंत्र के कार्य को भारतीय परिवेश, इसकी सामाजिक-आर्थिक स्थितियों, आकार, जनसंख्या तथा ऐतिहासिक पृष्ठभूमि के संदर्भ में देखा जाना चाहिए।

एक भाषा से दूसरी भाषा में अनुवाद, वह भी पूर्णतया भिन्न व्याकरणिक संरचनावाली भाषाओं का अनुवाद, के लिए बहुत व्यापक कलेवर (स्पैक्ट्रम) को समेटना पड़ता है। मंत्र ने स्पैक्ट्रम के एक भाग से शुरुआत कर दी है तथा यह स्थापित कर दिया है कि यह महान कार्य कुछ निर्दिष्ट उद्देश्यों की पूर्ति के लिए संपादित किया जा सकता है। यह लोगों तक पहुँचने के लिए रास्ता खोलता है, उन्हें शिक्षित करता है, उनके क्षितिज को व्यापक बनाता है तथा वैश्विक समाज में सक्रिय भूमिका निभाने योग्य बनाता है। इसके अनुप्रयोग समाज को इन लक्ष्यों के लिए तैयार करते हैं।

मशीनी अनुवाद के क्षेत्र में कई दशकों से कार्य किया जा रहा था तथा 90 के दशक के आरंभ में यह विश्वसनीय प्रौद्योगिकी कृत्रिम बुद्धि तथा संगणकीय भाषाविज्ञान के क्षेत्र में प्रगत शोधों के साथ उभरी। 1989 में सी-डेक के अनुप्रयुक्त कृत्रिम बुद्धि समूह द्वारा आरंभिक पहल की गई तथा आवर्धित ट्रांजिशन संजाल (ATN) तथा वृक्ष निकटस्थ व्याकरण (TAG) पदनिरूपक (Parser) का विकास किया गया। अंग्रेजी, हिंदी, गुजराती, संस्कृत तथा जर्मन में कार्य करने में सक्षम व्याकर्ता (VYAKARTA) नामक टैगपदव्याख्याता का निर्माण करने के बाद, संबंधित अनुप्रयोग की खोज की गई। भारतीय संदर्भ में अनुवाद अत्यावश्यक मामला था। अतः प्रथम वास्तविक जीवन अनुप्रयोग के रूप में केंद्र सरकार के विभागों में प्रयुक्त राजभाषा के क्षेत्र में अंग्रेजी-हिंदी के युगल का चयन किया गया। तदनुसार, आदिप्ररूप अनुवाद प्रणाली का सुविचारित निर्माण तथा उत्तरोत्तर परिष्कार तथा परिमार्जन करने के बाद “मंत्र” नाम से अभिमंत्रित किया गया।

मंत्र के इस संस्करण का प्रदर्शन राजभाषा विभाग (DOL) तथा अन्य संस्थाओं तथा संस्थानों के समक्ष किया गया। फलतः 1996 में राजभाषा विभाग (DOL) ने “प्रशासनिक उद्देश्यों के लिए कंप्यूटर की सहायता-प्राप्त अनुवाद प्रणाली” शीर्षक से एक परियोजना को प्रायोजित किया। इस उद्देश्य के लिए निर्दिष्ट क्षेत्र के रूप में सबसे पहले भारत सरकार में निगुक्ति की गजट अधिसूचना का चयन किया गया। इसके बाद कार्मिक प्रशासन में जारी अधिसूचनाओं के क्षेत्र बैंकिंग सैक्टर, विभिन्न मंत्रालय तथा अन्य एजेंसियों के कार्यक्षेत्र को भी शामिल किया जा रहा है।

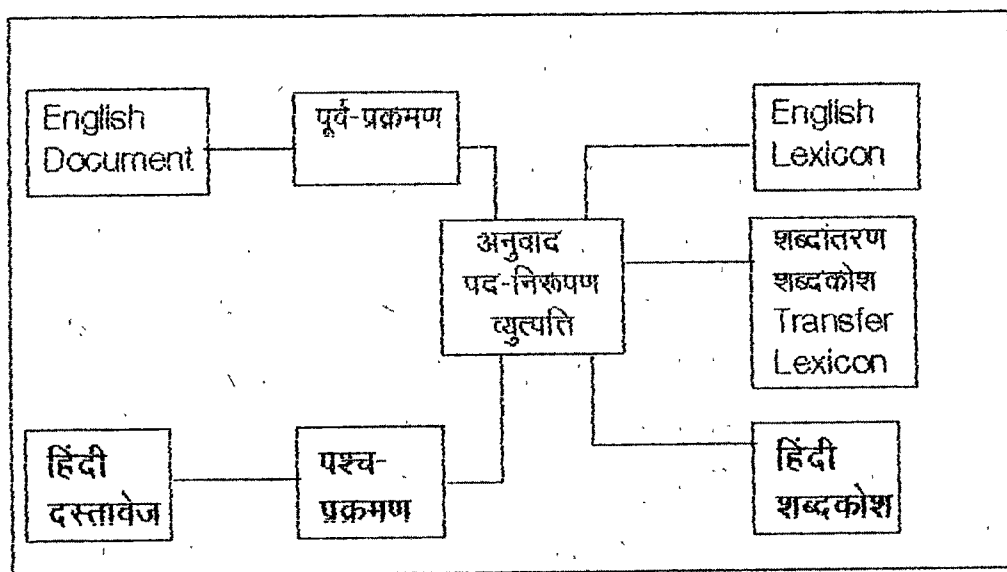
मंत्र का प्रयोग आरंभ में संस्थाओं तथा निगमों द्वारा किया जाना है, क्योंकि उन्हें कार्यालयीन प्रयोजनों के लिए काफी मात्रा में अनुवाद करना होता है। दीर्घकाल में भी यह किफायती होगा क्योंकि मनुष्य द्वारा किया जाने वाला अनुवाद कार्य का अधिकांश मशीनी अनुवाद प्रणाली द्वारा साझा कर लिया जाएगा। इसके प्रयोग के साथ इसका लाभ समुदाय को मिलेगा तथा अंततः भारतीय समाज का कोना-कोना यह स्वयं में समेट लेगा।

मंत्र कार्यनीति :

शब्द से शब्द तक नहीं ... नियम से नियम तक नहीं ... बल्कि शाब्दिक वृक्ष से शाब्दिक वृक्ष तक।

मंत्र का परिदृश्य

मंत्र में अंग्रेजी दस्तावेज को सबसे पहले पूर्व-प्रक्रमण (प्री-प्रोसेसिंग) हेतु भेजा जाता है। इसके बाद यह पदव्याख्या (पार्सिंग) तथा व्युत्पत्ति हेतु जाता है। इस दौरान अंग्रेजी शब्दकोश, हिंदी शब्दकोश तथा शब्दांतरण शब्दकोश का प्रयोग किया जाता है। व्युत्पत्ति-जनित हिंदी-आउटपुट तैयार करता है जिसे पश्च-प्रक्रमण (पोस्ट-प्रोसेसिंग) के लिए आवश्यकतानुसार भेजा जाता है तथा अंत में अनूदित हिंदी दस्तावेज प्राप्त होता है।



इनपुट/आउटपुट :

मंत्र सॉफ्टवेयर का इनपुट किसी भी मानक डीटीपी पैकेज या शब्द-संसाधक का प्रयोग करते हुए निर्मित समृद्ध पाठप्रारूप (RTF - Rich Text Format) फाइल या एचटीएमएल फाइल या प्लेन टेक्स्ट फाइल हो सकती है।

वाक-पहचान के क्षेत्र में विकास के साथ-साथ नई वाक-पहचान पैकेज जैसे IBM वाया-वोसज का बेहतर रूप सामने आ रहा है। इन पैकेजों का प्रयोग करते हुए माइक्रोसॉफ्ट वर्ड में डॉक्यूमेंट बना कर मंत्र में डाला जा सकता है। डॉक्यूमेंट को स्कैन भी किया जा सकता है, OCR प्रोग्राम चलाया जा सकता है तथा RTF फाइल बनाई जा सकती है जिसे अनुवाद हेतु सॉफ्टवेयर में स्थानांतरित किया जा सकता है। मंत्र की कड़ी माइक्रोसॉफ्ट वर्ड से सीधे जुड़ी हुई है, जो पैकेज के लिए डिफाल्ट वर्ड प्रोसेसर है। इनपुट अंग्रेजी पाठ का प्रारूपण (Formatting) अनूदित हिंदी आउटपुट में बरकरार रहता है। निर्दिष्ट ई-मेल पते पर संलग्नक के रूप में अनूदित हिंदी डॉक्यूमेंट भी भेजा जा सकता है।

प्रक्रमण के विभिन्न चरण

मंत्र प्रयोक्ता का मित्र है तथा सफलतापूर्वक प्रयोग में लाया जाने वाला सॉफ्टवेयर है। मंत्र में अनुवाद प्रक्रिया के तीन महत्वपूर्ण चरण हैं : (i) पूर्व-प्रक्रमण, (ii) पद-निरूपण (पार्सिंग) तथा व्युत्पत्ति, (iii) पश्च-प्रक्रमण।

(क) पूर्व-प्रक्रमण चरण

इस चरण के अंतर्गत अनेक कार्य सम्मिलित हैं जिसमें वर्तनी तथा व्याकरण की जाँच, संज्ञा, दिनांक तथा अन्य क्षेत्र-निर्दिष्ट-शब्दों की स्वतः पहचान, मैनुअल/स्वतः पदबंध निर्माण, पाठ में अंकन, वाक्य निष्कर्षण तथा शब्दकोश में खोज-कार्य शामिल है।

(ख) पद-निरूपण तथा व्युत्पत्ति

पद निरूपक गणन गहन मशीनरी है जो अंग्रेजी जो अंग्रेजी व्याकरण का प्रयोग करते हुए डॉक्यूमेंट का व्याकरण सम्मत तथा सीमैटिक (Semantic) विश्लेषण करता है। पदव्याख्याता अंग्रेजी वाक्य को व्युत्पत्ति-वृक्ष के रूप में पुनः प्रस्तुत करता है। व्युत्पत्ति-जनित्र इस व्युत्पत्ति-वृक्ष को आत्मसात कर के समतुल्य हिंदी वाक्य व्युत्पन्न करता है। यह कार्य वह अंग्रेजी-हिंदी स्थानांतरण शब्दकोश में उपलब्ध हिंदी व्याकरण तथा स्थानांतरण व्याकरण की सहायता से करता है। मंत्र के विकास में उल्लेखनीय मौलिक योगदान हिंदी टैग व्याकरण का विकास है।

(ग) पश्च-प्रक्रमण चरण

यह चरण वैकल्पिक है। इसके अंतर्गत पश्च-संपादन, अनेक हिंदी आउटपुट वाक्यों में से किसी एक वाक्य का चयन तथा आवश्यकतानुसार पश्च-प्रक्रमण यूटीलिटी जैसे प्रतिस्थापन, द्विभाषी आउटपुट तथा आउटपुट के परिमार्जन हेतु थिसारस/ शब्दकोश का प्रयोग शामिल है।

इस प्रकार मंत्र एक दर्शनबोध है... एक स्वप्न है... और अब वास्तविकता है।⁵

(ii) अनुवादक 2.0

अनुवादक 2.0 ऐसा सॉफ्टवेयर है जो अंग्रेजी मूल पाठ का हिंदी में तुरंत अनुवाद करने और हिंदी व्याकरण के मूलभूत नियमों का पालन करने में सक्षम है। कृत्रिम बुद्धि से युक्त यह सॉफ्टवेयर विभिन्न विषयों के अनुकूल अनुवाद कर सकता है। जैसे - सामान्य भाषाई अनुवाद, प्रशासनिक भाषा में अनुवाद, वैज्ञानिक भाषा में अनुवाद, कृषि के अनुकूल भाषा में अनुवाद आदि।

इस सॉफ्टवेयर में तीन अलग-अलग शब्दकोश भी अंतर्निहित हैं, जिनका इस्तेमाल अनुवादक के साथ या स्वतंत्र रूप से किया जा सकता है। पहला शब्दकोश सामान्य शब्दावली, दूसरा वैज्ञानिक शब्दावली व तीसरा कृषि संबंधी शब्दावली के लिए है। इनमें करीब डेढ़ लाख शब्द शामिल किए गए हैं।

अनुवादक छपे हुए पत्र में अंकित शब्दों को पहचानने में सक्षम है। किसी मुद्रित पृष्ठ को स्कैन कर के अनुवादक में फाइल के तौर पर सेव किया जाता है।

जिन अंग्रेजी शब्दों के समानार्थी हिंदी शब्द उपलब्ध नहीं हैं, उनके लिप्यंतरण की व्यवस्था भी अनुवादक में है। साथ ही, जिन अंग्रेजी शब्दों के एक से ज्यादा हिंदी अर्थ उपलब्ध हैं, उनके सभी संभव विकल्प प्रदर्शित किए जाते हैं ताकि शब्दों का सटीक चयन किया जा सके।

अनुवादक 2.0 ऐसा सॉफ्टवेयर समाधान है जो कृत्रिम बुद्धि से युक्त एक क्रांतिकारी अनुवाद इंजन की मदद से अंग्रेजी के दस्तावेजों का अनुवाद हिंदी भाषा और व्याकरण के नियमों का पालन करते हुए हिंदी में करता है। इस सॉफ्टवेयर में विंडोज के सभी मुख्य दस्तावेज रूपों के लिए परिवर्तन-फिल्टर उपलब्ध है और सॉफ्टवेयर के भीतर वर्ड-प्रोसेसर के भीतर से ही फाइलिंग और प्रिंटिंग की सुविधा के साथ-साथ आप अनुवाद से पहले और बाद, दोनों चरणों के लिए अंतर्निर्मित व्याकरण जाँच, ग्राמר चेकर द्वारा कर सकते हैं। सॉफ्टवेयर में विशेष कामों के लिए अंतर्समाहित शब्दकोशों से युक्त जैसे - औपचारिक, कृषि भाषा विज्ञान, तकनीकी और प्रशासनिक शब्दकोश व भव्य व कलात्मक 100 फॉन्ट तथा हिंदी कैरेक्टरों को अंग्रेजी कैरेक्टरों में मैप करने के लिए अंतर्निर्मित कैरेक्टर मैप हैं। साथ ही, शब्दकोश में नए शब्द जोड़ना या पुराने शब्दों में संशोधन करना बेहद आसान है। यदि पहले से विद्यमान किसी शब्द के अनेक अर्थ हैं तो उन्हें भी शब्दकोश में शामिल किया जा सकता है।

यहाँ सटीक अनुवाद के लिए जिन शब्दों के एकाधिक अर्थ उपलब्ध हैं उनमें से सबसे उपयुक्त शब्द चुनने की सुविधा उपलब्ध है तथा पुल-डाउन मेनू के कारण इस सॉफ्टवेयर को इस्तेमाल करना बहुत आसान है। साथ ही, सॉफ्टवेयर अंग्रेजी और हिंदी दोनों के लिए स्पेल चेकर के साथ अंग्रेजी मूल-पाठ और उसके हिंदी रूपांतरण दोनों को सेव करने की सुविधा प्रदान करता है।

सॉफ्टवेयर द्वारा प्रदान की जा रही सुविधाओं में रोजमर्रा के आधिकारिक पत्रों और सर्कुलरों के तीव्र, स्वतः अनुवाद का विकल्प और माइक्रोसॉफ्ट वर्ड के ही स्तर की संपादन

सुविधाएँ व अँग्रेजी के उन शब्दों के लिप्यंतरण की सुविधा है जिनका कोई हिंदी समानार्थक शब्द उपलब्ध नहीं है। सॉफ्टवेयर में एच.पी.स्टैंडर्ड स्कैनर के माध्यम से सेव की गई अँग्रेजी फाइलों को ला कर काम करने के लिए विभिन्न प्लग-इन उपलब्ध हैं तथा पॉवर पाइंट स्लाइड, पेंट-ब्रश में बने चित्रों और माइक्रोसॉफ्ट वर्ड के दस्तावेजों पर भी काम करना संभव है।

अनुवादक 2.0 सार्वजनिक एवं निजी सेक्टर, बैंको, प्रकाशकों, समाचारपत्रों और वकीलों के लिए आदर्श रूप में उपयोगी है जहाँ वार्षिक रिपोर्टों, समाचारपत्र-पत्रिकाओं,

पुस्तकों, बड़ी संख्याओं में परिपत्रों, अधिसूचनाओं, पत्राचार, बैंकिंग आदि में बड़ी मात्रा में अनुवाद किया जाता है।⁶

(iii) नामानुवाद (N-Trans)

यह ऐसा सॉफ्टवेयर है जो अँग्रेजी से भारतीय भाषाओं में तथा इसके विपरीत भारतीय भाषाओं से अँग्रेजी भाषा में लिप्यंतरण करता है। एन-ट्रांस शब्दकोश तथा अन्वेषण (heuristic)-सुझाव-मैकेनिज्म का विद्वतापूर्ण समन्वय है। इसके पहले भाग में शब्दकोश है जिसमें अत्यंत महत्वपूर्ण नामों को मूलधार रूप में गरा गया है। इसका दूसरा अवयव शक्तिशाली अन्वेषण (heuristic)-सुझाव-मैकेनिज्म है जो पाठ-संचालित है। यह मैकेनिज्म मनुष्य के समान बनने की प्रतिस्पर्धा करता है तथा यह बताने का प्रयास करता है कि किस शब्द का प्रयोग सर्वाधिक उपयुक्त रहेगा। एक शक्तिशाली इंजन द्वारा संचालित सांख्यिक वेटेज नियमों तथा पर्यावरण नियमों की सहायता से यह मानव-व्यवहार की सादृश्यता करते हुए उसकी नकल करता है तथा यथासंभव सर्वाधिक उपयुक्त सुझाव देता है।

यह अधिकांशतः सही वर्तनी उपलब्ध कराता है। गलत वर्तनी अथवा अप्रचलित वर्तनियों के कारण नियमों का पूरा अनुपालन न होने की स्थिति में एन-ट्रांस सर्वाधिक उपयुक्त ध्वन्यात्मक वर्तनी के अनुसार काम करता है तथा कई मामलों में सही प्रमाणित होता है।

एन-ट्रांस में व्यक्तिगत शब्दकोश तैयार करने की सुविधा उपलब्ध है जिसमें विशेष नामों तथा शब्दों को शामिल किया जा सकता है।

प्रमुख विशेषताएँ :

- पाठ तथा डाटाबेस फाइलों के रूपांतरण हेतु उपयोगिता के रूप में उपलब्ध।
- मॉड्यूल द्वारा जब और जैसी आवश्यकता होने पर अनुवाद हेतु किसी अनुप्रयोग में अंतःसंबद्ध किए जाने के लिए पुस्तकालय के रूप में भी उपलब्ध।
- नियमित रूप से भिन्न वर्तनियों के शब्दों हेतु व्यक्तिगत शब्दकोश बनाने (अपेडिंग) का विकल्प।
- 85% तक सही अनुवाद।
- यह दक्षिण भारतीय भाषाओं के साथ-साथ उत्तर भारतीय भाषाओं के लिए भी उपलब्ध है।

कुछ अनुप्रयोग :

दूरभाष डाइरेक्टरी, वेतनपंजी (पे-रोल), बीजक (इनवोइसिंग), कस्टम अनुप्रयोग आदि।⁷

(iv) लिप्यंतरण (Transliteration) इसके तहत एक भाषा के पाठ को दूसरी भाषा में बदल दिया जाता है। ध्वन्यात्मक भाषाओं के संदर्भ में यह काफी उपयोगी होता है। आई लीप में यह सुविधा उपलब्ध है।

"यदि आप ट्रांसलिट्रेशन का कार्य करना चाहते हैं तो आपको आईलीप के टूल मीनू में ट्रांसलिट्रेट कमांड मिलेगा, जो एक भाषा के टेक्स्ट को दूसरी भाषा में परिवर्तित कर देगा। परिवर्तन करने पर यहाँ पर टेक्स्ट का अर्थ न बदल कर केवल अक्षर बदलते हैं, जिससे दूसरी भाषा के लोग उसको पढ़ कर भी न समझ सकें। जैसे ही ट्रांसलिट्रेट नामक यह फंक्शन क्रियान्वित होता है, इसका ऑप्शन मीनू सामने आ जाएगा। इसमें सबसे पहले स्क्रिप्ट का चुनाव करें और फिर ओ.के. बटन पर क्लिक कर दें। ऐसा करने से चुनी हुई स्क्रिप्ट के अनुसार टेक्स्ट बदल जाएगा। लेकिन आप जिस टेक्स्ट को बदलना चाहते हैं, उसका सलेक्ट होना आवश्यक है अन्यथा यह फंक्शन काम नहीं करेगा।"⁸

(v) शब्दानुवाद (Word Translation)

इसके तहत किसी भाषा के पाठ का अनुवाद एक-एक शब्द का अनुवाद करते हुए किया जाता है। इसके बाद प्रयोक्ता को उन शब्दों को लक्ष्य भाषा के व्याकरण के अनुसार स्वयं व्यवस्थित करना पड़ता है क्योंकि पूरे-पूरे वाक्यों का अनुवाद गड़बड़मिटा मिलता है। लीप ऑफिस यह सुविधा उपलब्ध कराता है।

(vi) पर्यायवाची शब्दकोश (Synonym Dictionary)

इसके तहत किसी शब्द के अनेक पर्यायों को उपलब्ध कराया जाता है जिससे वांछित शब्द का प्रयोग किया जा सके। यह सुविधा लीप ऑफिस में उपलब्ध है।

(vii) राजभाषा हिंदी शब्दकोश (Official Language Hindi Dictionary)

"यह प्रशासनिक शब्दों तथा पदबंधों का ऑनलाइन शब्दकोश है। यह सामान्यतः प्रयुक्त प्रशासनिक शब्दों तथा पदबंधों को उपलब्ध कराता है तथा अनेक भारतीय भाषाओं के लिए भी सहायक है।"⁹

आईएसएम-2000 तथा लीप ऑफिस में यह सुविधा उपलब्ध है।

(viii) शक्ति - त्रिभाषी अनुवादक (SHAKTI - Trilingual Translator)

"भारत में कंप्यूटर और इंटरनेट उपयोक्ताओं की संख्या बढ़ाने के लिए चल रही कोशिशों के तहत अंतरराष्ट्रीय सूचना प्रौद्योगिकी संस्थान (इंटरनेशनल इंस्टीट्यूट ऑफ इंफॉर्मेशन टेक्नोलॉजी, आई आई आई टी), हैदराबाद ने कई नए समाधान प्रस्तुत किए हैं। संस्थान ने ग्रामीण इलाकों के उपयोक्ताओं की जरूरतों का ध्यान रखते हुए सारे समाधान विकसित किए हैं। आई आई आई टी, हैदराबाद भारतीय भाषाओं के अनुवाद का एप्लीकेशन (अनुप्रयोग) तैयार कर रहा है। पहले दौर में संस्थान ने जो एप्लीकेशन तैयार

MANTRA

Dear Sir,

I reached Pune on Wednesday, 27th March. I am in the Electronics Department of the Foundation Day Ceremony. Dr. C. N. Rao will deliver the Foundation Day Lecture. I will be in the CDAC Pune for six days.

Dr. M. Chandrasekhar has taken over as Director, CDAC Bangalore. I shall be in his office on 27th March, 1999. I will return to Pune on 28th March, 1999.

Yours sincerely,
Rajesh Kumar Singh

English Source Text

Dear Sir,

I reached Pune on Wednesday, 27th March. I am in the Electronics Department of the Foundation Day Ceremony. Dr. C. N. Rao will deliver the Foundation Day Lecture. I will be in the CDAC Pune for six days.

Dr. M. Chandrasekhar has taken over as Director, CDAC Bangalore. I shall be in his office on 27th March, 1999. I will return to Pune on 28th March, 1999.

Yours sincerely,
Rajesh Kumar Singh

Body Text Selection

Dear Sir,

I reached Pune on Wednesday, 27th March. I am in the Electronics Department of the Foundation Day Ceremony. Dr. C. N. Rao will deliver the Foundation Day Lecture. I will be in the CDAC Pune for six days.

Dr. M. Chandrasekhar has taken over as Director, CDAC Bangalore. I shall be in his office on 27th March, 1999. I will return to Pune on 28th March, 1999.

Yours sincerely,
Rajesh Kumar Singh

English Text Parsing

Dear Sir,

I reached Pune on Wednesday, 27th March. I am in the Electronics Department of the Foundation Day Ceremony. Dr. C. N. Rao will deliver the Foundation Day Lecture. I will be in the CDAC Pune for six days.

Dr. M. Chandrasekhar has taken over as Director, CDAC Bangalore. I shall be in his office on 27th March, 1999. I will return to Pune on 28th March, 1999.

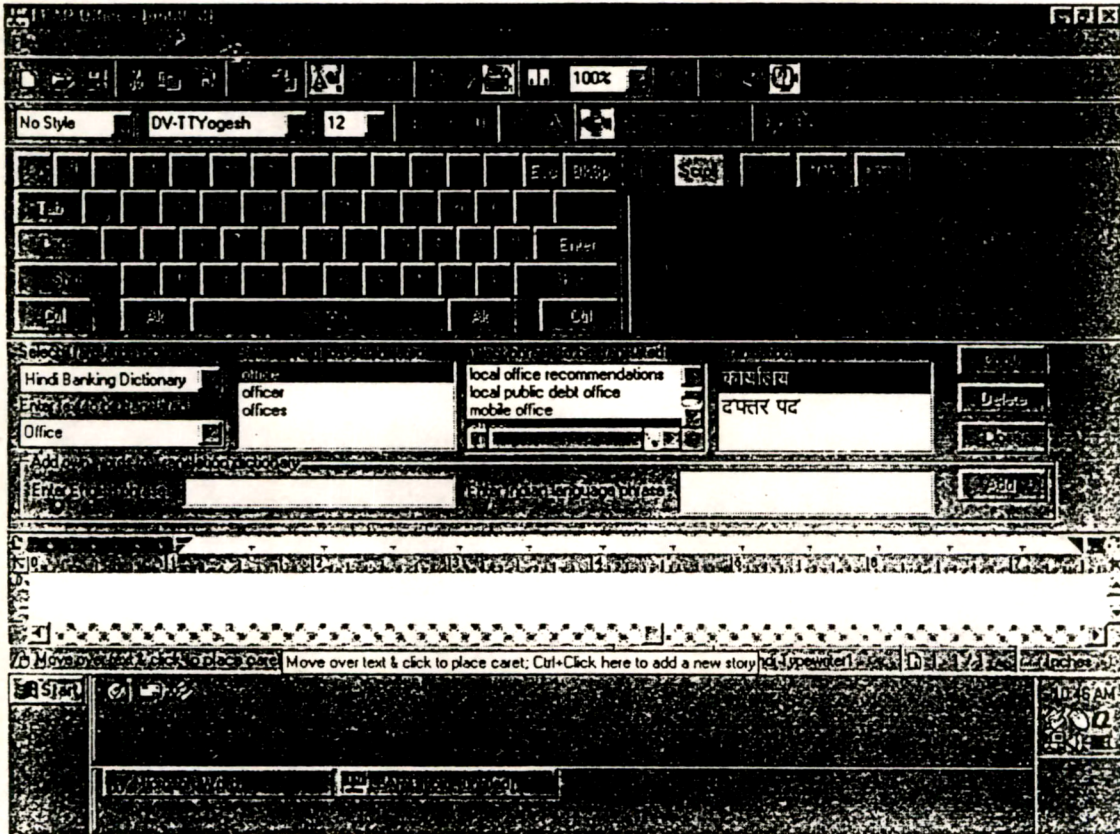
Yours sincerely,
Rajesh Kumar Singh

English-Hindi Translated Output

लीप ऑफिस

अंग्रेजी से भारतीय भाषाओं में वर्ड ट्रांसलेशन

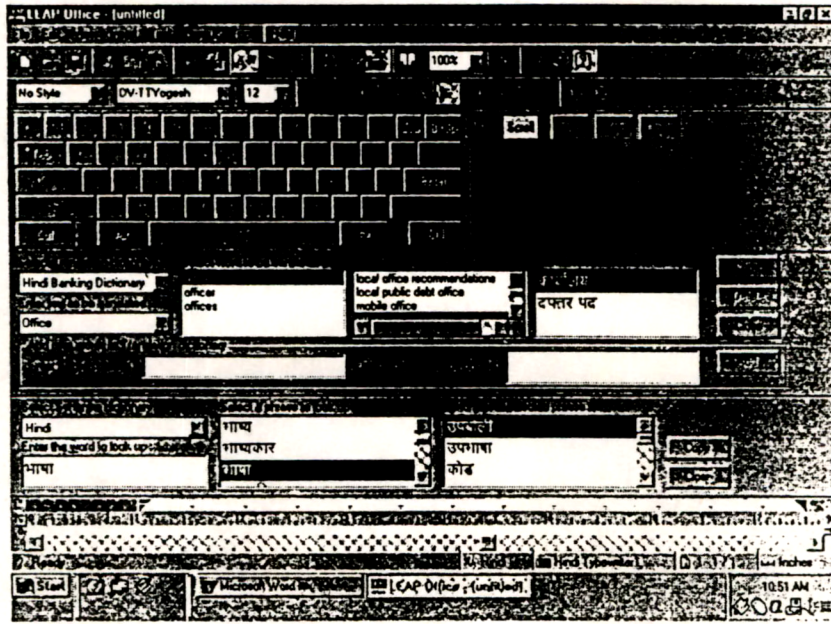
लीप ऑफिस में Hindi Administrative Dictionary तथा Hindi Banking Dictionary उपलब्ध है।



लीप ऑफिस

पर्यायवाची शब्द Synonym Dictionary

Tools >> Synonym Dictionary



Enter English phrase:

□ 附 錄

Hindi Administrative Dictionary

Account

Select The Word

Select The Phrase

Translated Words

account
accountability
accountable
accountancy
accountant

- account
- account book
- account head
- annual account
- appropriation account

लेखा
हिसाब
खाता

Copy To Clipboard

Cancel

किया है, वह अँग्रेजी टेक्स्ट (पाठ) का दो भारतीय भाषाओं - हिंदी और तेलुगु में अनुवाद कर सकता है। इसके बाद इसे दूसरी भाषाओं में भी अनुवाद करने लायक बनाया जाएगा। इस एप्लीकेशन का कोड नाम शक्ति रखा गया है। इसे अमेरिका के कार्नेगी मेलान यूनीवर्सिटी के साथ तालमेल कर (सहयोग से) तैयार किया गया है। इसे तैयार करने वाली टीम के प्रमुख राजीव संगल का मानना है कि इस एप्लिकेशन के इस्तेमाल से कोई भी भारतीय भाषा को जानने वाला कंप्यूटर-उपयोक्ता अँग्रेजी की बातों को अपनी भाषा में देख और समझ सकता है।¹⁰

6.4 आवाज पहचानक/वाक् अभिज्ञान (SPEECH/VOICE RECOGNITION)

6.4.1 विकास गाथा

स्पीच रेकॉग्निशन तकनीक कोई नई नहीं है। कंप्यूटर वैज्ञानिक पिछले चालीस साल से इसे विकसित करने में लगे हैं। लेकिन कंप्यूटर आपकी बात सुने और समझे, यह अभी तक संभव नहीं हो सका है और इस काम को अभी लंबा रास्ता तय करना है। स्पीच रेकॉग्निशन तकनीक के विकास के प्रमुख चरण इस प्रकार हैं :

1950 का उत्तरार्द्ध : आईबीएम के वैज्ञानिकों ने विशाल कंप्यूटरों के उपयोग से स्पीच रेकॉग्निशन पर शोध शुरू किया। इन वैज्ञानिकों ने ध्वनि और शब्दों के बीच सांख्यिकीय सह संबंध तलाशने का काम शुरू किया।

1960 का दशक : अमेरिकी रक्षा विभाग ने कुछ स्पीच रेकॉग्निशन शोध के लिए कोष बनाया। वैज्ञानिकों ने 'ऑटोमेटिक प्रोटोटाइपिंग' विकसित किया। यह एक ऐसी प्रक्रिया है जिसमें कंप्यूटर विशिष्ट ध्वनियों का पता लगाकर उन्हें संग्रहित करता है।

1964 : आर्थर सी क्लार्क और स्टेनले क्यूब्रिक ने '2001-ए स्पेस ओडिसी' में सहयोग दिया। फिल्म का HAL-9000 कंप्यूटर, पूर्णरूप से फंक्शनल इंटरएक्टिव कंप्यूटर था।

1982 : ड्रैगन सिस्टम विकसित किया गया। दो साल बाद इंग्लैंड में इसकी पहली स्पीच रेकॉग्निशन तकनीक को पोर्टेबल पीसी में लगाया गया।

1993 : आईबीएम ने पहला स्पीच रेकॉग्निशन उत्पाद पैकेज पेश किया। ओएस/2 के लिए आईबीएम पर्सनल डिक्टेशन सिस्टम था। मैकिंटोश के लिए एप्पल ने अपना उत्पाद उतारा।

1994 : ड्रैगन सिस्टम ने पहला सॉफ्टवेयर आधारित डिक्टेशन उत्पाद पेश किया। यह विंडोज-1.0 के लिए ड्रैगन डिक्टेट था।

1995 : ओएस/2 में स्पीच नेविगेशन और रेकॉग्निशन सिस्टम भी साथ आए।

1996 : आईबीएम ने पहला वास्तविक स्पीच रेकॉग्निशन उत्पाद, मैड स्पीक/रेडियोलॉजी पेश किया।

1997 : ड्रैगन ने पहला सामान्य उपयोग वाला स्पीच रेकॉग्निशन उत्पाद पेश किया, आईबीएम ने वाया वॉइस उतारा और माइक्रोसॉफ्ट ने लर्नराइट एंड हॉस्पी पेश किया।

1998 : बाजार में नेचुरल लैंग्वेज प्रोग्रामों की बाढ़-सी आ गई।

1999 : शोधकर्ता अब रोजमर्रा के अनुप्रयोगों के लिए और नए स्पीच व वॉइस रेकॉग्निशन उत्पाद तैयार करने में जुटे हैं।

6.4.2 आवाज पहचानक : कैसे और क्या ?

जब हम बोलते हैं तो ध्वनि उत्पन्न होती है और यह ध्वनि कंपन हवा के माध्यम से समांतर तरंग (एनालॉग वेव) के रूप में चलती है। इस एनालॉग वेव में मूल्यों (वैल्यू) के असंख्य समूह होते हैं। जब एनालॉग वेव का आयाम बढ़ता और कम होता है, तो ध्वनि के सुर में भी उसी हिसाब से वृद्धि और कमी होती है, यही ध्वनि का सबसे वास्तविक रूप होता है। आजकल के कंप्यूटर, एनालॉग इनपुट की जगह डिजिटल सूचनाओं को संचालित करते हैं। ध्वनि की व्याख्या कर आउटपुट देने के लिए कंप्यूटर को एनालॉग तरंगों को डिजिटल डेटाओं में तब्दील करना होता है। एनालॉग वेव को डिजिटल डेटाओं में बदलने की प्रक्रिया ही **सैंपलिंग** कहलाती है। ध्वनि की संपूर्ण एनालॉग वेव को रिकार्ड करने की बजाय, कंप्यूटर के साउंड कार्ड में एनालॉग वेव, डिजिटल में बदल जाती है और एक ही समय में विभिन्न बिंदुओं पर तरंग की स्थिति के शॉट आ जाते हैं। क्वांटीकरण के जरिए इन बिंदुओं को एक मान (वैल्यू) दे दिया जाता है, जिन्हें बाद में बाइनरी कोड की डिजिटल भाषा में परिवर्तित किया जाता है, जिसे कि कंप्यूटर समझता है। जब साउंड कार्ड ध्वनि के सैंपल रखता है, तो यह उसके मान को रिकार्ड कर लेता है। जब ध्वनि को डिजिटल रूप में परिवर्तित किया जाता है, तो परिणाम स्टेयर्स के सेट के रूप में ज्यादा मिलता है, बजाय रैंप के। वह आवृत्ति जिस पर साउंड स्नैप शॉट लिए जाते हैं, उसे **सैंपलिंग रेट** कहा जाता है।

करीब साढ़े तीन दशक पहले आर्थर सी क्लार्क और स्टेनले क्यूब्रिक ने '2001 ए स्पेस ओडिसी' में कंप्यूटर और मानव के बीच संवाद-संभव की परिकल्पना को पेश किया था। तब एचएएल-9000 कंप्यूटर का बयान इस प्रकार था, 'मुझे खेद है डेव मैं डरा हुआ हूं, मैं यह नहीं कर सकता।' लेकिन कंप्यूटर (एचएएल-9000) के इस कथन ने दुनिया भर के उन लोगों के मन में ढेरों विचार जरूर पैदा कर दिए थे जो आदमी और कंप्यूटर के बीच बातचीत और आपसी समझ बनने का सपना देख रहे थे।

कंप्यूटर और मानव के बीच आपसी समझ और संवाद को अभी तक हम '2001' और 'स्टार ट्रेक' सीरिज में ही देख पाए हैं। यह कृत्रिम बौद्धिक विज्ञापन का एक ऐसा स्तर है जिसे अभी हमें हासिल करना है। इन कृत्रिम बौद्धिक तरीकों में कंप्यूटर पर बात

करने की बजाय कंप्यूटर के साथ बातचीत के तरीके शामिल हैं। उदाहरण के लिए आप अपने कंप्यूटर को किसी ब्रह्मांडी पिंड का रास्ता तलाशने को कहते हैं और दूसरी ओर उससे बोल कर एक चिट्ठी लिखवाते हैं - तो इन दोनों में काफी फर्क है। यह फर्क ठीक वैसा ही है जैसे कंप्यूटर 'पर' बात करने और कंप्यूटर 'के साथ' बात करने में है।

आज जो स्पीच रेकॉग्निशन सॉफ्टवेयर आ रहे हैं, उनकी मदद से आप अपने कंप्यूटर से साफ-साफ तेज बात कह सकते हैं। हालांकि कंप्यूटर यह नहीं समझ पाएगा कि आप क्या कह रहे हैं। वह केवल सुन सकता है, व्याख्या कर सकता है और जवाब दे सकता है। लेकिन उसे ऐसा अंतर्ज्ञान नहीं है कि वह आपकी बात को समझ ले।

आजकल ज्यादातर लोग ऐसे फोनो का इस्तेमाल करते हैं जिसमें स्पीच रेकॉग्निशन तकनीक पहले ही से होती है। कई बैंक और दूसरे सेवा संस्थान इंटरएक्टिव वॉइस रिसर्गोस (आईवीआर) सिस्टम से सेवाएं देते हैं।

स्पीच रेकॉग्निशन सिस्टम और वॉइस रेकॉग्निशन में फर्क है। स्पीच रेकॉग्निशन सॉफ्टवेयर जो आप वास्तव में कहते हैं उसका जवाब देता है। यह सॉफ्टवेयर शब्दों की पहचान करता है और आपको उचित रूप से जवाब देता है, जबकि वॉइस रेकॉग्निशन सॉफ्टवेयर सुरक्षा उद्देश्यों के लिए काफी ज्यादा उपयोग होता है।

स्पीच रेकॉग्निशन बोलने के साथ ही शुरू होती है। हम सब यह मानकर चलते हैं कि यह हर एक कोई जानता है कि विचार तंत्रिकातंत्र में ही उत्पन्न होते हैं और किस प्रकार मस्तिष्क इन विद्युतीय मनोवेगों को मांसपेशियों की गति में परिवर्तित करता है और मांसपेशियों की गति ही आवाज के रूप में सामने आती है जिससे कि हम आपस में बात कर पाते हैं।

आवाज या किसी भी प्रकार की ध्वनि, तरंग के रूप में माध्यम में ही गति करती है। माध्यम हवा या पानी अथवा निर्वात हो सकता है। जब ये तरंगें किसी दूसरी वस्तु से टकराती हैं तो कंपन उत्पन्न होती है। इसी तरह जब तरंगें हमारे कानों से टकराती हैं तो संकेतों में बदल जाती हैं और ये संकेत हमारे मस्तिष्क तक पहुंचते हैं। फिर मस्तिष्क इन संकेतों को ध्वनियों में बदल देता है। दूसरी ओर कंप्यूटरों में माइक्रोफोन का प्रयोग किया जाता है जो इलेक्ट्रॉनिक कान की तरह काम करते हैं।

माइक्रोफोन कई तरह के आते हैं, लेकिन इन सभी में एक तरह की पतली सतह होती है जो ध्वनि तरंगें टकराने पर कंपन शुरू करती हैं। इसी से थोड़ा-सा विद्युत प्रवाह होता है। फिर सतह पर कंपन उत्पन्न होती है जो विद्युत संकेतों में समान विभिन्नता पैदा करती है। ये संकेत फिर आपके कंप्यूटर के साउंड कार्ड को एक जैक के जरिए भेजे जाते हैं, जहाँ माइक्रोफोन मशीन से जुड़ा होता है।

कंप्यूटर के भीतर सबसे पहला काम यह होता है कि इन संकेतों को इस प्रकार रूपांतरित किया जाता है ताकि पीसी इन्हें समझ सके। अगर हम एनालॉग संकेत का कागज पर ग्राफ उतारें तो यह एक तरंग की तरह ही नजर आएगा। और यह तरंग की आवृत्ति या तीव्रता को प्रदर्शित करेगा।

कंप्यूटर डिजिटल सूचनाओं से संचालित होते हैं। ये डिजिटल सूचनाएं अंकों से प्रदर्शित होती हैं। ये सूचनाएं यह नहीं बताती हैं कि संकेत कैसे बदल रहे हैं, बल्कि ये बताती हैं कि एक समय में किसी विशिष्ट बिंदु पर ये कैसे आती हैं अथवा प्रदर्शित होती हैं। एनालॉग संकेत (सिग्नल) को डिजिटल संकेत में बदलने के लिए सैंपलिंग की जरूरत होती है। सैंपलिंग की प्रक्रिया में एनालॉग संकेत को दो पतली परतों में बांट दिया जाता है। उदाहरण के लिए मान लीजिए कि पहले नैनोसैकेंड से दूसरे नैनोसैकेंड के बीच संकेत (सिग्नल) का मान (वैल्यू) '1' से बदल कर '3' हो जाता है। सैंपलिंग की सहायता से इन नैनोसैकेंड सेक्शन को औसत 'दो' तक किया जा सकता है और तब इसे शून्य और एक के बाइनरी कोड में लिखा जाता है। इसी तरह हरेक खंड एक मान प्राप्त करता है। तकनीकी रूप से एक एनालॉग सिग्नल जितना परिशुद्ध होता है, उतना डिजिटल सिग्नल नहीं।

जैसे ही ध्वनि, अंकों के कॉलम में रूपांतरित होती है, तभी सॉफ्टवेयर का असली काम शुरू हो जाता है। सॉफ्टवेयर यह पहचानता है कि ये सब अंक किन शब्दों को बता रहे हैं। हमें लगेगा कि एक से शब्द या एक-सी ध्वनि हमेशा अंकों के एक ही समूह पर मिलें। तभी यह कहा जा सकेगा कि विशिष्ट नंबर पर ही विशिष्ट ध्वनि है। वास्तव में व्यवहार में यह एक जटिल समस्या है। प्रोग्रामरों को स्पीच रेकॉग्निशन सॉफ्टवेयर विकसित करते समय अलग-अलग लोगों की भाषाई विशेषता को ध्यान में रखना चाहिए। अगर आप किसी व्यक्ति को ध्यान से सुनेंगे, तभी आपको यह पता लग जाएगा कि उसकी बोली की खासियत क्या है और यही वह असली जानकारी होती है, जो स्पीच रेकॉग्निशन सॉफ्टवेयर डिजाइन करने में मदद देती है। अलग-अलग व्यक्तियों और क्षेत्रों की अपनी-अपनी बोलियां होती हैं, उनकी अपनी ध्वनि होती है, अलग शैली होती है। कंप्यूटर डिजिटल कोड तो यह सब पहचानेगा कि गलत सूखा है, बोलने वाला थका हुआ-सा है, व्यक्ति धीरे बोल रहा है या तेज। यह सारा काम काफी जटिल और समस्या भरा होता है। किसी भी वाक्य को समझने के लिए जो सबसे पहली जरूरत होती है वह यह कि हम उन अधिकतर ध्वनियों को समझें जिनसे कि वह वाक्य बना है। भाषा के इन खंडों को वैज्ञानिक फोनीम कहते हैं।

स्पीच रेकॉग्निशन का काम वाक्य ध्वनि को छोटे-छोटे बराबर के हिस्सों में बांटने से शुरू होता है। इन्हें वैक्टर्स कहते हैं। वैक्टर्स में से किसी भी फोनीम को उठाने से पहले सॉफ्टवेयर डेटाओं को सामान्य कर देता है ताकि शोर न्यूनतम हो जाए और दूसरे अंतर भी न्यूनतम स्तर तक आ जाएं। यह पूरी प्रक्रिया एक गणितीय विधि 'लघुगणक' (अलॉगरिथम) के जरिए ही संपन्न होती है।

अब इन सामान्यीकृत वैक्टर्स पर सॉफ्टवेयर अपना काम शुरू करता है। इस स्थिति में भी कंप्यूटर पूरी तरह फोनीम के सही जोड़े नहीं बना सकता है। सॉफ्टवेयर के पास ढेर सारी डेटा सूचनाएं होती हैं जिनकी जरूरत फोनीमों को पड़ती है। यदि ये यह पता नहीं लगा पाते हैं कि विशिष्ट ध्वनि कौन-सी है, तब सॉफ्टवेयर इनकी मदद कर सकता है।

सही फोनीम भी उपयुक्त शब्द नहीं बना पाते हैं। जैसे 't' और 'ooh' की ध्वनि को एक साथ रखें तो जो शब्द बन सकते हैं, वे ये होंगे - 'to' या 'two' या फिर 'टू'।

यदि आप कहते हैं - आई, तो क्या आपका सही मतलब 'eye' से होता है ? इसलिए वॉइस रेकॉग्निशन तकनीक में संदर्भ प्रमुख होता है।

ट्राइग्राम एनालिसिस : इसमें तीन-तीन शब्दों का समूह होता है। जिनमें दो शब्द तीसरे उपयुक्त शब्द से मिलकर सार्थक वाक्य बनाते हैं। इस प्रकार तीन-तीन शब्दों के तमाम समूहों को लेकर कंपनियों ने अपने सॉफ्टवेयर बनाए हैं। आईबीएम ऐसा सॉफ्टवेयर लाई है जिसमें अलग-अलग भाषाओं के दो शब्द तीसरे विशिष्ट शब्द से मिलेंगे। तीन-तीन शब्दों का यह समूह ही एक बड़ा डेटाबेस तैयार करता है। पिछले दो शब्दों से मिलने वाले तीसरे शब्दों की संख्या को शोधकर्ता ब्रॉचिंग फैक्टर कहते हैं। उदाहरण के तौर पर एक्स-रे के काम में लगे रेडियोलॉजिस्ट के लिए यह ब्रॉचिंग फैक्टर बीस दफ्तरी कामकाज में 150 और अखबार मैगजीनों के लिए तीन सौ तक होता है।

ट्रेनिंग सॉफ्टवेयर : विभिन्न बोलियों और भाषाओं से इकट्ठे किए गए आंकड़ों का उपयोग करके सॉफ्टवेयर निर्माता ऐसे प्रोग्राम तैयार कर सकते हैं, जो कि उन ज्यादातर शब्दों की पहचान कर सकते हैं जो बोले जाते हैं। आईबीएम ने इस काम में काफी हद तक सफलता पाई है और उसके 95 फीसदी शब्दों का चयन इस प्रकार है जो रोजमर्रा में बोले-समझे जाते हैं, इसलिए इनमें गलती होने की संभावना बहुत ही कम है।

बायोमैट्रिक्स : बायोमैट्रिक्स अनुप्रयोगों में उपयोग किए जाने वाले वॉइस रेकॉग्निशन की प्रक्रिया भी स्पीच रेकॉग्निशन तकनीक पर ही आधारित है। हालांकि इनमें एक फर्क यह है कि स्पीच रेकॉग्निशन में कंप्यूटर का सारा ध्यान इस पर केंद्रित होता है कि बोलने वाला वाक्य कैसे बोलता है, जबकि वॉइस रेकॉग्निशन हर आदमी के बोलने के तरीके और बोली की खास विशेषताओं पर आधारित होती है। ज्यादातर बायोमैट्रिक अनुप्रयोग में यह जरूरी होता है कि आप विशिष्ट शब्दों को बोलें। इसके बाद आपकी आवाज की कुछ निश्चित विशिष्टताओं का विश्लेषण करता है - जैसे ध्वनि, स्वरसंक्रम आदि। इसके बाद सिस्टम हर उपयोक्ता के लिए सांचों (टैम्पलेट) का निर्माण करता है और इन सांचों में डेटाबेस संग्रह करता है।

बायोमैट्रिक सिस्टम पहचान या पुष्टि के लिए ही संचालित किए जाते हैं। पहचान में यह जरूरी होता है कि व्यक्ति सिस्टम को अपनी आवाज का नमूना दर्ज कराए, ताकि सांचे में संग्रहीत ध्वनि से इसका मिलान किया जा सके। जबकि वैरिफिकेशन (पुष्टि) में व्यक्ति के बारे में जानकारी देने वाला प्रारंभिक संकेत जैसे नाम या दिया हुआ विशिष्ट नंबर भी दिया जाना जरूरी होता है। तीन प्रमुख क्षेत्र ऐसे हैं जिनमें स्पीच या वॉइस रेकॉग्निशन तकनीक का इस्तेमाल मुख्य रूप से किया जाता है। ये हैं - डिक्टेशन, इंटरफेस और सिक्योरिटी।¹¹

6.4.3 विभिन्न आवाज पहचानक/वाक् अभिज्ञान

(क) नेचुरली स्पीकिंग का इंटरफेस वायावॉइस के इंटरफेस से थोड़ा-सा ज्यादा सहज ज्ञानी है। इसका परिशुद्धता-केंद्र भी उपयोग करने में ज्यादा आसान है, जो टूलबारों और मेनुओं से पूरी तरह युक्त है। इसमें आप पहले ही से डिक्टेड कराई गई ध्वनि को आसानी से सुन सकते हैं। इसकी ध्वनि आईबीएम की रोबोटिक ध्वनि से कहीं ज्यादा सहज है।

किसी भी तरह के पार्श्व शोर को हटाने के लिए ड्रैगन का हैंडसेट सॉफ्टवेयर के साथ अच्छी तरह काम करता है और लंबे व अनावश्यक पॉज हटा देता है। दूसरी ओर वायावॉइस में सबसे बड़ी खराबी यह है कि अगर आप बुदबुदा भी रहे हैं तो 'इट', 'अम' जैसे छोटे अनावश्यक शब्दों को इंसर्ट कर लेता है। वैसे विशिष्टता के तौर पर देखें तो वायावॉइस में की-कंट्रोल की सुविधा है जिससे आप माइक्रोफोन चालू कर सकते हैं और रिकार्डिंग हो जाने पर उसे तत्काल रोक भी सकते हैं।

(ख) नेट कंट्रोलर : जब आप वेब सर्फिंग पर हों तो वायावॉइस काफी काम का है। इसे विंडोज एक्सपी के साथ काम करने के लिए डिजाइन किया गया है। यह शार्टकटों और वेब एड्रेसों को फौरन पकड़ लेता है। इसकी कंट्रोल कमांडें जैसे 'डिलीट दिस पैराग्राफ' एकदम दोबारा काम करती हैं। सूचनाओं के अपने भंडार से सेक्शन की मदद करती हैं।

(ग) एक्सपी में स्पीच रेकॉग्निशन : माइक्रोसॉफ्ट ने अपने ऑफिस एक्सपी में स्पीच रेकॉग्निशन सुविधाएं दी हैं। इस प्रकार उपयोगकर्ताओं के लिए यह एक आदर्श सोल्यूशन के रूप में प्रतिष्ठित हुआ है। जिन लोगों को कभी-कभी की-बोर्ड की मदद लेने की जरूरत पड़ सकती है, उनके लिए यह लाभदायक है, लेकिन इसे वॉइस रेकॉग्निशन पैकेज के पावर की जरूरत नहीं है।

डॉक्यूमेंट डिक्टेशन के लिए यह एक संपूर्ण रूप से स्वीकार्य सॉफ्टवेयर है। हालांकि आप पाएंगे कि माइक्रोसॉफ्ट के पैकेज में की-बोर्ड का उपयोग कहीं ज्यादा करना पड़ जाता है क्योंकि इसमें डिक्टेशन और कमांड मोड के बीच आपको मैनुअल स्विच की जरूरत पड़ती है। पीछे से आने वाले शोर की समस्या इसमें भी है। हल्की बड़बड़ाहट की ध्वनि भी इसमें आ जाती है।

(घ) ड्रैगन नेचुरली स्पीकिंग - 6.0 : 400 मेगाहर्ट्ज पेंटियम-II, विंडोज 98/Me/2000/एक्सपी, 128 एम बी रैम, 300 एमबी हार्ड डिस्क स्पेस, यूएसबी पोर्ट (माइक्रोफोन उपयोग के लिए), 16 बिट का साउंड कार्ड।

(ङ) वायावॉइस-10.0 : 300 मेगाहर्ट्ज पेंटियम, 256 L2 कैश, विंडोज-98 के लिए 64 एमबी रैम, विंडोज एक्सपी के लिए 128 एमबी, 510 एमबी हार्ड डिस्क स्पेस, यूएसबी पोर्ट, 16 बिट साउंड कार्ड।

(च) ऑफिस एक्सपी स्टैंडर्ड : 400 मेगाहर्ट्ज पेंटियम-II, 128 एमबी रैम, यूएसबी माइक्रोफोन, ऑपरेटिंग सिस्टम इंस्टाल करने के लिए सामान्य जरूरतें।

6.4.4 विभिन्न आवाज पहचानक उत्पादों का विश्लेषण

नए वॉइस रेकॉग्निशन सॉफ्टवेयर पैकेज आए हैं, उनके नतीजे शत-प्रतिशत शुद्ध हैं। इनकी कीमत भी काफी कम है।

माइक्रोसॉफ्ट ने वायदा किया था कि उसका टैबलेट पीसी तकनीकी से रुबरु होने के हमारे तरीके को बदल डालेगा। यह टैबलेट पीसी बिना की-बोर्ड का होगा। इसमें की-बोर्ड की जगह ज्यादा सहज पेपर और पैन सिस्टम होगा। लेकिन वाक-शक्ति ही होगी जो हमें डेस्क से पूरी तरह अलग कर देगी।

दूसरी तकनीकियों में जिस तेजी से सुधार और विकास हुआ है, वहीं स्पीच रेकॉग्निशन को अपने को बनाए रखने के लिए काफी मुश्किलें झेलनी पड़ी हैं। इसका कारण अपर्याप्त प्रोसेसिंग पावर और अंग्रेजी भाषा की जटिलता है। लेकिन अब मल्टी-गीगाहर्ट्ज स्पीड वाली मशीनें आने से और कम दाम के पैकेजों से शुद्धता की दर काफी बढ़ गई है।

वॉइस रेकॉग्निशन एप्लिकेशन आज हमारे चारों ओर मौजूद हैं। ऑटोमेटेड टेलीफोन सेवाओं से आपके मोबाइल तक वॉइस आती है, जिससे आप बिना डायलपैड के किसी व्यक्ति से संपर्क कर सकते हैं।

वॉइस रेकॉग्निशन पर भाषाएं सिखाने वाले प्रोग्राम भी हैं। ये प्रोग्राम उपयोक्ताओं को उच्चारण करना तक सिखाते हैं, ताकि वे भाषण के बारे में अच्छी जानकारी

हासिल कर सकें और बोलना सीख सकें। लेकिन इस तरह की तकनीक को विकसित करना काफी महंगा है। इस क्षेत्र में मुख्य रूप से दो ही बड़ी कंपनियाँ हैं - आईबीएम और स्कैनसॉफ्ट। आईबीएम का वायावॉइस है और स्कैनसॉफ्ट का ड्रैगन नेचुरली स्पीकिंग सॉफ्टवेयर है। वॉइस रेकॉग्निशन सॉफ्टवेयर के लिए आईबीएम ने 1970 के दशक से ही काम शुरू कर दिया था। इसका पहला उत्पाद 1990 के आसपास बाजार में आया जो वॉइस रेकॉग्निशन सॉफ्टवेयर से तैयार किया गया था। लेकिन 1997 तक इस दिशा में ज्यादा कुछ नहीं हुआ। इसी साल अपना वायावॉइस (VIAVOICE) ब्रांड जारी किया।

आईबीएम के साथ ही प्रतिस्पर्धा में आई स्कैनसॉफ्ट ने 'लर्नआउट एंड हॉस्पी' खरीद कर इस पर काम शुरू किया। 'लर्नआउट एंड हॉस्पी' पहले 'ड्रैगन' शीर्षक से था, जो कई खामियों की वजह से चल नहीं पाया था। दिसंबर, 2001 में स्कैनसॉफ्ट ने इसे लिया और इसकी स्पीच फैसिलिटी सुधारने पर काम शुरू किया।

नतीजतन इन पैकेजों में शुद्धता 97 फीसदी तक आई और इनकी बोली (वाक क्षमता) भी कहीं ज्यादा सहज बनी। आज जैसी रोबोटिक स्पीच संभव है, उससे कहीं

ज्यादा बेहतर ये सॉफ्टवेयर पैकेज हैं। फिर भी आखिर ऐसी क्या वजह है कि अभी भी बड़ी तादाद में लोग स्पीच रेकॉग्निशन सॉफ्टवेयर पसंद नहीं कर रहे हैं ? इसका एक कारण यह है कि वॉइस रेकॉग्निशन सॉफ्टवेयर के प्रशिक्षण और अभ्यास में समय लगता है। दूसरा कारण यह है कि इस सॉफ्टवेयर का उपयोग एकदम शांत माहौल में होना चाहिए। दफ्तरों में आम तौर पर शोर-शराबा होता ही है जो इनके लिए उपयुक्त जगह नहीं हो सकती।

अब माइक्रोसॉफ्ट भी इस दिशा में पीछे नहीं है। माइक्रोसॉफ्ट अपने ऑफिस एक्सपी सूइट के जरिए स्पीच रेकॉग्निशन सिस्टम को मुख्यधारा में ला रहा है। अगर उपयोक्ताओं के पास टूल हैं, तो वे इसका इस्तेमाल कर सकते हैं। ऑफिस एक्सपी के दोनों वर्जनों पीसी और टैबलेट वर्जन में वॉइस रेकॉग्निशन क्षमताएं हैं।

ज्यादातर उपयोक्ताओं को यह जानकारी नहीं है कि ऑफिस एक्सपी में वॉइस रेकॉग्निशन सुविधा है। कंपनी इसे खुलेआम स्वीकार करती है कि वॉइस रेकॉग्निशन एक सेकंडरी एप्लिकेशन ही बना रहेगा क्योंकि ज्यादातर उपयोक्ता डिजिटल इंक और हैंडराइटिंग रेकॉग्निशन (हस्तलेखन पहचान) को ज्यादा प्राथमिकता दे रहे हैं।

ऑफिस एक्सपी की स्पीच फैसिलिटी को व्यापक तौर पर नहीं अपनाए जाने का एक जो बड़ा कारण है, वह अमेरिकी-अंग्रेजी के वॉइस रेकॉग्निशन इंजन का होना है। ब्रिटेन के अंग्रेजों के लिए यह बड़ी समस्या पैदा करता है। अमेरिका बाजार में बढ़ती मांग के कारण स्पीच रेकॉग्नाइजर अमेरिकी-अंग्रेजी में तैयार किए गए हैं। आने वाले समय में इसके दूसरे इंटरनेशनल वर्जन तैयार होंगे।

कैसे काम करता है स्पीच रेकॉग्निशन सॉफ्टवेयर :

स्पीच रेकॉग्निशन सॉफ्टवेयर बोले गए शब्दों को कमांडों में परिवर्तित कर देता है। फिर ये कमांड पीसी से व्याख्या के रूप में परिवर्तित होते हैं। कंप्यूटर यह नहीं समझता है कि आप क्या बोल रहे हैं। जबकि आप जो कुछ बोलते हैं, माइक्रोफोन उसे एनालॉग सिग्नल में परिवर्तित कर देता है। फिर ये सिग्नल साउंड कार्ड में जाते हैं। इसलिए आपके पीसी में जितना अच्छा साउंड कार्ड होगा, शुद्धता की दर भी उतनी ही ज्यादा होगी।

इसके बाद एडीसी (एनालॉग टू डिजिटल कंवर्टर) इन सिग्नलों को डिजिटल डेटाओं की स्ट्रीम में परिवर्तित कर देता है जिन्हें फोनम कहते हैं।

यहां तक आने पर ध्वनि साइल किसी भी इंटरफियरेंस (व्यतिकरण) को हटा देता है, जैसे- बैकग्राउंड शोर। इसके बाद यह डेटाओं को पिचों व साउंड में समेट देता है, उनका विश्लेषण करता है और फिर शब्दों को फोनम के डिजिटल प्रदर्शन में परिवर्तित कर देता है।

इसके बाद लैंग्वेज-माडल डिजिटल प्रदर्शनों को वापस शब्दों में बदलता है। यह काम इंटरनल-डिक्शनरी को यादृच्छिक रूप से चेक करके करता है। हालांकि इसमें समस्या यह है कि इसमें कई शब्दों की ध्वनि एक-सी होती है, लेकिन अर्थ उनके अलग-अलग होते हैं। ऐसी स्थिति में प्रोग्राम ट्राइग्राम का उपयोग करता है। इसमें दो-दो शब्दों को समूह के रूप

में पहले या बाद में रखा जाता है और फिर दिए गए संदर्भ में कौन-सा शब्द सही बैठता है, यह पता लगा कर उसका उपयोग किया जाता है।

6.4.5 प्रशिक्षण का महत्व

इन कामों को दक्षतापूर्वक निपटाने के लिए यह जरूरी है कि सॉफ्टवेयर उपयोक्ता का स्पीच पैटर्न पहचानता हो। संदर्भ का पता लगाने के लिए यह पूर्व में बनाए गए वाक्यों का उपयोग करेगा। इसलिए जितने ज्यादा डेटा संग्रहित होंगे, परिशुद्धता उतनी ही ज्यादा आएगी। लेकिन 10 से 15 मिनट में इसे उपयोगकर्ता की आवाज को पकड़ लेना चाहिए।

ये दोनों ही पैकेज ऐसे हैं जो इंस्टाल करते समय उपयोक्ताओं को पढ़ने के लिए प्रोत्साहित करते हैं। हालांकि नेचुरली स्पीकिंग (NATURALLY SPEAKING) शुरू करने के लिए बेहतर है और वायावॉइस थोड़ा आपके धैर्य का परीक्षण ले सकता है। यदि यह किसी शब्द को पहचान नहीं पाता है तो यह उसे हाईलाइट कर देगा, टेक्स्ट रोक देगा और फिर जब तक जवाब नहीं मिल जाता तब तक संकेत ध्वनि देता रहेगा।

नेचुरली स्पीकिंग की एप्रोच ज्यादा बढ़िया है। उपयोक्ता को लगातार पढ़ने के लिए प्रोत्साहित करता है चाहे सॉफ्टवेयर उसमें एक भी शब्द छोड़े या नहीं छोड़े। इसके बाद यह उस शब्द को उठाकर अगली बार वाले पैसेज में रख देता है। हालांकि इसमें थोड़ा-सा गतिरोध महसूस होता है, लेकिन प्रक्रिया ज्यादा सहज लगती है।

वायावॉइस 10.0 के साथ आईबीएम की पेशकश ज्यादा अच्छी है। यह सॉफ्टवेयर तेज और इंस्टाल करने में आसान है। हालांकि इसकी शुरुआत में परिशुद्धता की दर कम है। यह नए शब्दों को तेजी से उठा लेता है।

नेचुरली स्पीकिंग-6.0 के कई वर्जन उपलब्ध हैं। लोटस नोट्स के साथ लगाने के लिए बनाए गए वर्जन से लेकर कई व्यावसायिक एप्लिकेशन तक उपलब्ध हैं। ये सारे रोजाना काम आने वाले एप्लिकेशनों के लिए ही तैयार किए गए हैं। जैसे - वर्ड, वर्डपरफेक्ट और एक्सेल।

वैरी घरेलू उपयोक्ताओं के लिए छोटा वर्जन सबसे ठीक है। इससे आप वर्ड में डिकटेड कर सकते हैं और इंटरनेट पर छानबीन कर सकते हैं, लेकिन इसके लिए आपको एक हैंडसेट या प्रीस्टैंडिंग माइक्रोफोन खरीदना होगा।

आईबीएम का वायावॉइस 10.0 तीन वर्जनों में आता है। कारोबारियों के लिए प्रो-यूएसबी वर्जन ठीक है। आईबीएम ने भी कारोबारियों और घरेलू उपयोक्ताओं के लिए कम कीमत का विंडोज स्टैंडर्ड एडिशन बनाया है।

वॉइस रेकॉग्निशन उन लोगों के लिए उपयोगी होता है जो बीमार हों, विकलांग हों। यह उन लोगों के लिए उपयुक्त नहीं है जिन्हें नजर नहीं आता है। अगर उन लोगों के हाथ टाइप करने में सक्षम हैं तो उनके लिए यह ठीक है। जो लोग हाथ या बाजूओं से अक्षम हैं उनके लिए यह एक विकट समस्या बन जाता है। स्पीच रेकॉग्निशन सॉफ्टवेयर

अभी पूरी तरह से उस मुकाम पर नहीं पहुंच पाया है, जहां विकलांग लोग इसे एक कमांड सेंटर के रूप में उपयोग कर सकें।

दरअसल लागत और पीसी नहीं ले पाने से संबंधित समस्याओं के कारण दृष्टिहीनों में तकनीकी साक्षरता काफी कम है, इसलिए वॉइस रेकॉग्निशन तकनीक उन विकसित हो रही तकनीकियों में से एक है, जो भविष्य में एक बड़ी भूमिका अदा करेगी।

6.4.6 भविष्य

वॉइस रेकॉग्निशन का भविष्य काफी उज्ज्वल है। कार निर्माता कंपनियां- होंडा व बीएमडब्ल्यू ने इसे ड्राइविंग के भविष्य के रूप में देखना शुरू कर दिया है। आपको सिर्फ बटन दबाने की जरूरत होगी। आप स्टॉरियो को नियंत्रित कर सकेंगे, बोल कर ही हीटिंग सिस्टम और लाइटों का नियंत्रण संभव होगा।

मोबाइल फोन कंपनियों ने तो पहले ही से वॉइस रेकॉग्निशन के उपयोग की संभावनाओं को टटोल लिया था। अब तो कई घरेलू इस्तेमाल के उपकरण भी ऐसे आ रहे हैं जिन्हें आप बोल कर नियंत्रित कर सकते हैं। आने वाले समय में तो टैबलेट पीसी और दूसरी सभी हाथ की डिवाइस वॉइस-रेकॉग्निशन-सॉफ्टवेयर की सुविधा से युक्त होंगी। नेचुरली स्पीकिंग मोबाइल वर्जन वॉइस-मेमो-रिकार्डर के साथ आता है, जो कि एक बिजनेस एप्लिकेशन है, लेकिन वॉइस रेकॉग्निशन को अभी पूरी तरह मुख्यधारा में लाना बाकी है। इन डिवाइसों के लिए मुख्य डेटा इनपुट विधि के रूप में ज्यादा-से-ज्यादा लोग वॉइस-रेकॉग्निशन तकनीक का उपयोग करेंगे।

पीडीए (पर्सनल डिजिटल असिस्टेंस) में आईबीएम में वायावॉइस रेकॉग्निशन सॉफ्टवेयर का उपयोग तेजी से बढ़ा है। इसके इंटरफेस इस तरह डिजाइन किए गए हैं कि उपयोग करने के लिए काफी आसान रहें, लेकिन सॉफ्टवेयर को चालू करने में जो समय लगता है, उससे कई बार झुंझलाहट होती है।¹²

6.5 सोनोग्राफ

सोनोग्राफ नामक यंत्र ध्वनि संकेतों का विश्लेषण कर सकता है। मशीन द्वारा किए गए अध्ययन के फलस्वरूप प्राप्त बर्हिवेश ध्वनि खंड (सोनोग्राफ) कहलाता है। यह बोले गए शब्दों की पहचान करने में सहायक होता है और उनके अर्थ निकाल सकता है। यह अब भी लगातार एक चुनौती बना हुआ है और उसके अनुसंधान पर अब तक काफी धनराशि खर्च हो चुकी है। मशीनों से बोलना मानव-मशीन संवाद का सबसे प्राकृतिक और सक्षम तरीका होगा।

आँखें और हाथ दोनों ही व्यस्त हों (जैसे कि हवाई जहाज के पायलट के) तब ऐसी स्थिति में बोलना ही संचार का सबसे प्राकृतिक ढंग है। अपने आप ध्वनि की पहचान करने वाले, ऑटोमैटिक स्पीच रिकॉग्नाइजर (ASR) ऐसी इलेक्ट्रॉनिक युक्तियाँ हैं जो हमारे द्वारा बोले गए शब्दों को ग्रहण करती हैं और उससे प्राप्त बर्हिवेश, कोडित डिजिटल संकेतों के रूप में होता है। इन्हें अनुप्रयोग एकक (एप्लीकेशन यूनिट) में जैसे का तैसा भरा जा सकता है। इन संकेतों से ही मालूम होता है कि ASR ने स्वर संकेत से क्या समझा है। ASR से

इन संकेतों को ग्रहण करने के बाद अनुप्रयोग एकक इस पर उचित कार्यवाही करता है। वे एक विशेष मशीनी क्रिया भी प्रारंभ कर सकते हैं। उदाहरण के लिए, एक रोबोट को चलने के, घूमने के, और भी न जाने कितने निर्देश मुँह जुबानी दिए जा सकते हैं। व्यापारिक रूप से उपलब्ध ASRs में स्वर ध्वनियों की ऐसी कोई बोधगम्यता नहीं है। बहिर्वेश संकेत अभिज्ञात शब्दों के बिल्कुल अनुरूप होते हैं।

मनुष्य द्वारा उच्चारित शब्दों का या वाणी का स्रोत, ध्वनि से पहचाना जाता है। मुख के जिस भाग से कोई वर्ण बोला जाता है, वही उस वर्ण का उच्चारण स्थान माना जाता है। ये स्थान हैं : कंठ, तालु, मूर्ध्ना, दंत, ओष्ठ तथा नासिका। इन स्थानों से बोले जाने वाले वर्ण क्रमशः कंठ्य (क, ख, ग), तालव्य (च, छ, ज, झ), मूर्धन्य (ट, ठ, ड, ढ, ण), दंत्य (त, थ, द, ध, न), ओष्ठ्य (प, फ, ब, भ, म) और नासिक्य (उ, ऋ, ए आदि) कहलाते हैं। किंतु कुछ वर्णों का उच्चारण दो-दो स्थानों से भी होता है जैसे व (दंत्योष्ठ्य), ओ, औ (कंठोष्ठ्य) और उ (कंठ्य एवं नासिक्य) आदि।

स्वर ध्वनियों के भाषाई अध्ययन से मालूम हुआ है कि किसी भी भाषा में मूल अमूर्त ध्वनि खंड होते हैं और ये बोलते समय आपस में मिल कर शब्दों की रचना करते हैं। इन खंडों को 'ध्वनिखंड' कहते हैं। अधिकांश बोली जाने वाली भाषाओं के औसतन 40 ध्वनिखंड होते हैं। इनमें से प्रत्येक ध्वनिखंड को विशिष्ट गुणों से पहचाना जाता है।

तथापि कंप्यूटर तंत्रों द्वारा ध्वनिखंडों की पहचान में अभी तक सफलता नहीं मिली है। अलग-अलग ध्वनिखंडों पर लगातार स्वर का प्रासंगिक प्रभाव इतना यथेष्ट होता है कि मशीन के लिए पूरी निश्चितता के साथ उसे समझ पाना कठिन होता है। पुरुषों और स्त्रियों के बोलने के ढंग में बहुत अंतर होता है। एक वक्ता द्वारा विभिन्न अवसरों पर बोले गए एक ही शब्द में भी कुछ अंतर हो सकता है। वही शब्द दूसरे वक्ताओं द्वारा बोले जाने में भी अंतर हो सकता है। ये प्रत्येक व्यक्ति के बोलने के लहजे और ढंग पर निर्भर करता है। पहली स्थिति में अंतर, शब्द में ध्वनिखंड की स्थिति और संदर्भ के कारण होता है। दूसरी स्थिति में, यह अंतर स्वर यंत्र के आकार और प्रत्येक व्यक्ति के बोलने के ढंग आदि कारकों के कारण होता है। एक वक्ता से पूर्णतया स्वतंत्र ध्वनि-अभिज्ञान-तंत्र के विकास के लिए, ध्वनि संकेतों में होने वाले अंतर के कारणों की पूरी समझ की भी जरूरत है।

सिलसिलेवार शब्दों में लगभग 300 मिलीसेकेंड के अंतराल के बाद अलग-अलग शब्दों की पहचान करने वाले उपकरण बाजार में उपलब्ध हैं। ये मशीनें विशेष रूप से, एक जाने पहचाने वक्ता के अलग-अलग शब्दों के छोटे-छोटे सैट की पहचान करती हैं। अधिकांश तंत्र बैंड-पास फिल्टर्स के एक बैंक की सहायता से प्रचलित स्पैक्ट्रम बनाते हैं। इन मशीनों के शब्दकोष में 20 और 200 के बीच ध्वानिक रूप से भिन्न शब्द होते हैं। प्रत्येक 20 मिलीसेकेंड के बाद शब्दों के प्रत्येक स्पैक्ट्रम की जाँच की जाती है। विभिन्न आवृत्ति के बैंडों के ऊर्जा अंश की तुलना स्मृति में पहले से ही उपलब्ध और संग्रहित संवादी बैंडों से की जाती है और निवेश से सबसे ज्यादा मिलते जुलते टुकड़े की पहचान करके उसे शब्द के रूप में चुन लिया जाता है।

शब्दों की पहचान करने वाले अलग-अलग तंत्रों में 'डायनेमिक प्रोग्रामिंग' नामक तकनीक का बहुत प्रयोग किया जाता है। यह हल खोजने वाला एल्गोरिदम है। दो एक जैसे प्रतिदर्शों के बीच बेहतर पहचान करने के लिए इसका प्रयोग किया जाता है। एक ही वक्ता द्वारा अलग-अलग अवसरों पर बोले गए शब्दों में अलग-अलग अंतराल हो सकता है इसलिए, ^{समय}संरेखण आवश्यक होता है। अगर तंत्र बड़ी शब्दावली का प्रबंध करता है तो एक मैच प्रदर्शित करने के लिए समान अनुपात में न केवल अधिक समय की जरूरत होगी बल्कि हर नए वक्ता के लिए अधिक प्रशिक्षण काल की भी जरूरत होगी। इस सब में बहुत समय लगेगा।

शब्दों के बीच मनचाहा अंतराल देकर, सिलसिलेवार ध्वनि संकेतों को अलग-अलग शब्दों को अनुक्रम में सीमित किया जा सकता है। इस स्थिति में, अलग-अलग शब्दों को पहचाना जा सकता है। यह विधि बोले गए शब्दों को स्क्रीन पर देखने के लिए पर्याप्त है या जहाँ एकल शब्द निर्देशों से बनी सूचनाओं की पहचान करनी हो। इसे 'प्रशिक्षण काल' कहते हैं। प्रत्येक उपशोक्ता के पास, मशीन का उपयोग करने से पहले अपनी बोली के सांचे होने चाहिए। तब ही वह मशीन का उपयोग कर सकता है। हालांकि ऐसे तंत्रों में जहाँ वक्ता कोई भी हो, अलग-अलग प्रशिक्षण योजनाओं की जरूरत नहीं पड़ती। ऐसे तंत्रों में आमतौर पर बहुत से वक्ताओं के सांचों से औसत सांचा बनाया जाता है और उसे मशीन की स्मृति में सहेज दिया जाता है। ऐसे वक्ताओं पर निर्भर न होकर शब्दों की पहचान करने वाले सांचे, शब्दों के ध्वनिक स्वरूप पर आधारित सूचनाओं पर बने होते हैं, उनके विस्तृत स्पेक्ट्रम के आधार पर नहीं।

कंटीन्युअस स्पीच रेकॉग्निशन (CSR) का कार्यक्षेत्र अत्यंत व्यापक है। प्राकृतिक स्वर में अनेक उतार चढ़ाव होते हैं। सहउच्चारण (अर्थात् एक शब्द का आगे आने वाले शब्द पर प्रभाव) और लय के कारण ध्वनि की लगातार पहचान करना कठिन हो जाता है। ग्रहण करने योग्य वाक्यों की सीमा की पहचान के लिए, पहचान करने वाले यंत्र हमारा भाषा - विज्ञान संबंधी ज्ञान का प्रयोग (व्याकरण सम्मत और अर्थ विषयक दोनों) करते हैं। इनके ध्वनीय और अक्षरीय प्रतिरूपों दोनों में ही बड़ी-बड़ी शब्दावलियां होती हैं। तकनीकी के इस क्षेत्र में पहले से विस्तृत अनुसंधान जारी है, CSR जिनका मुख्य लक्ष्य है। एक प्रतिशत त्रुटि दर वाले तंत्र पहले ही बनाए जा चुके हैं। हालांकि अलग-अलग मूल्यांकनों में त्रुटि दर 12 प्रतिशत और 0.5 प्रतिशत के बीच देखी गई है। ये तंत्रों की जटिलता और उनके मूल्य पर निर्भर करती है।

जब ध्वनि संकेतों में, भेजे जा रहे संदेश से संबंधित भाषा विषयक सूचना भी होती है, तो उसमें भाषा के अलावा कुछ अतिरिक्त सूचनाएँ भी होती हैं जैसे वक्ता की पहचान, लहजा, उसकी मनोवैज्ञानिक एवं क्रियात्मक स्थिति और इनके साथ-साथ उस समय की पारिवेशिक स्थितियाँ जैसे शोर, कक्ष ध्वनिकी आदि। उच्च स्तरीय ध्वनि अभिज्ञान तंत्र बनाने के लिए किसी को भी पहले यह सीखना होगा कि संकेत से संदेश-धारी घटक को कैसे अलग किया जाता है और बाकी को छोड़ दिया जाता है। स्पष्ट है कि अभी ध्वनि विज्ञान में अत्यंत मौलिक अनुसंधान की आवश्यकता है तब कहीं ये कृत्रिम ध्वनि अभिज्ञान तंत्र मनुष्य की क्रियाओं के आसपास पहुँच सकेंगे।¹³

6.6 वृक्ष-संलग्न व्याकरण (टैग)

TREE ADJOINING GRAMMAR (TAG)

भाषा वाक्यों का समुच्चय (Set) है। प्रत्येक वाक्य भाषा-विशेष की शब्दावली के प्रतीकों (शब्दों) की संरचित माला (Structured String) है। भाषाविज्ञान में भाषा विशेष की उसी संरचना का निरूपण कतिपय नियमों के आधार पर किया जाता है। इन्हीं नियमों के समुच्चय को व्याकरण कहा जाता है। प्रजनक व्याकरण (Generative Grammar) के अंतर्गत इन नियमों के रूप (Form) और प्रकार्य (Functioning) अलग-अलग होते हैं।

प्राकृतिक भाषा को कंप्यूटर के माध्यम से समझने, विश्लेषित और संश्लेषित करने के लिए अनेक प्रविधियों का उपयोग किया जाता है। किंतु इन सभी प्रविधियों में सबसे महत्वपूर्ण बात है प्राकृतिक भाषा में प्रयुक्त वाक्यों का अर्थ निर्धारण।

प्राकृतिक भाषा संसाधन प्रणाली के अंतर्गत व्याकरण एक महत्वपूर्ण उपादान (Crucial Component) है। यही कारण है कि अधिकांश विशेषज्ञों ने ऐसे ही रूपवाद विकसित करने का प्रयास किया है, जो प्राकृतिक भाषाओं का व्याकरण लिखने के लिए उपयुक्त सिद्ध हों। इस प्रकार के अधिकांश रूपवाद संदर्भमुक्त व्याकरण (Context Free Grammar - CFG) के ही विस्तार हैं। व्याकरण-रूपवाद से संबंधित प्रमुख समस्या है, वाक्य-विज्ञान को अर्थ-विज्ञान से किस स्तर पर अलग किया जाए। वाक्यात्मक विश्लेषण में किन भाषिक तत्वों को समाहित किया जाए और अर्थपरक विश्लेषण के लिए क्या छोड़ा जाए ?

संदर्भमुक्त व्याकरण वह व्याकरण है जो भाषा की संरचना का वर्णन अत्यंत सरल रूप में करता है और संचयित ज्ञान (Stored Knowledge) वाक्यों के निर्माण (Production) और अभिज्ञान (Recognition) के लिए उत्पन्न संरचना और निर्माण तथा अभिज्ञान की प्रक्रिया के बीच संबंध स्थापित करता है। इसे परंपरागत भाषा-वैज्ञानिक तात्कालिक अवयव व्याकरण (Immediate Constituent Grammar) के नाम से जानते हैं। वस्तुतः यह एक विशेष प्रकार का पदबंध संरचना व्याकरण (Phrase Structure Grammar) ही है। पदबंध-संरचना-व्याकरण प्रजनक भाषा वैज्ञानिकों और प्राकृतिक भाषाओं या प्रोग्रामन भाषाओं के प्रकलन (Manipulating) के लिए अपेक्षित कंप्यूटर प्रणालियों का मुख्य आधार है।

कुछ परिस्थितियों में हम संदर्भ पर कुछ प्रतिबंध (Constraints) लगाना चाहेंगे, तब उसे संदर्भयुक्त व्याकरण (Context Sensitive Grammar - CSG) कहा जाएगा। संदर्भयुक्त व्याकरण आवर्ती (Recursive) होता है। इसमें कभी-कभी वैकल्पिक संकेतों (Alternative Notions) का प्रयोग किया जाता है। इस वैकल्पिक संकेत प्रणाली में इस बात पर जोर रहता है कि प्रतीक का पुनर्लेखन संदर्भ पर निर्भर है।

संदर्भमुक्त और संदर्भयुक्त दोनों प्रकार के व्याकरणों की सीमाएँ हैं। जहाँ एक ओर संदर्भमुक्त व्याकरण में आवर्ती नियमों की व्यवस्था नहीं है, वहाँ संदर्भयुक्त व्याकरण में चयनात्मक प्रतिबंध केवल बाह्य संरचना तक ही सीमित रहते हैं, जबकि प्राकृतिक भाषाओं में अर्थपरक विश्लेषण के बिना अनेक वाक्यों का पदनिरूपण सही नहीं होगा।

• 6.6.1 संदर्भमुक्त व्याकरण (CFG) तथा टैग (TAG) में अंतर

वृक्ष संलग्न व्याकरण अर्थात् टैग (TAG) संदर्भमुक्त-व्याकरण के विकास की अगली कड़ी है। CFG में प्रत्येक निगम केवल यही निर्देश करता है कि पदनिरूपण-वृक्ष के ठीक आगे के ताल का निर्माण किस प्रकार किया जाए। यदि हम वृक्ष की गहराई को और बढ़ाएंगे, तो भी व्याकरण की शक्ति बढ़ेगी नहीं क्योंकि हमें केवल इस बात की अनुमति है कि हम वृक्ष की पत्तियों पर प्रतिस्थापना का कार्य करें। यदि किसी रूप में यह संभव हो जाए कि वृक्ष के आंतरिक नोड पर भी कुछ संलग्न किया जा सके तो स्वाभाविक है कि उस रूपवाद (Formalism) की शक्ति स्वतः ही बढ़ जाएगी। टैग में ठीक यही शक्ति को संलग्नक (Adjoining) कहा जाता है।

TAG और CFG में दूसरा अंतर यह है कि जहाँ CFG शब्द-गुच्छ प्रजनन प्रणाली (String Generating System) है, तो वहाँ टैग वृक्ष प्रजनन प्रणाली (Tree Generating System) है।

संलग्नक नाम की संयोजन क्रिया में और अधिक शब्द-गुच्छों (Strings) को व्युत्पन्न (Derive) करने की गुंजाइश रहती है। संलग्नक प्रतिस्थापना की तुलना में कहीं अधिक शक्तिशाली है तथा प्रतिस्थापना को उत्प्रेरित कर सकता है। इसकी सहायता से उन भाषाओं को भी प्रजनित किया जा सकता है, जिनका प्रजनन प्रतिस्थापना की सहायता से संभव नहीं था। प्रतिस्थापना और संलग्नक की सहायता से ही CFG को शब्दीकृत किया जाता है।

6.6.2 टैग का रूपवाद (Formalism)

सन 1973 में Tree Adjoining Grammar (TAG) शीर्षक से अरविंद जोशी, लेवी और ताकाहासी का लेख प्रकाशित होने के बाद उक्त व्याकरण को एक रूपवाद (Formalism) के रूप में मान्यता मिली। वृक्ष संलग्न व्याकरण (TAG) में जो प्रजनक शक्ति है, वह संदर्भमुक्त व्याकरण (CFG) और संदर्भयुक्त व्याकरण (CSG) के कहीं बीच में निहित है। संभवतः इसीलिए इसे हल्का-सा संदर्भयुक्त व्याकरण (Mildly Context Sensitive Grammar) कहा जाता है। इसके अंतर्गत दो प्रकार की वृक्ष संरचनाएँ होती हैं : आदि वृक्ष (Initial Tree) और गौण वृक्ष (Auxiliary Tree)। इन दोनों वृक्षों को समवेत रूप में मूल वृक्ष (Elementary Tree) कहा जाता है।

आदि वृक्ष न्यूनतम वाक्य वृक्ष है और गौण वृक्ष "X" लेबल के रूप में चिह्नित मूल और पाद नोड सहित न्यूनतम आवर्ती संरचना है। दोनों ही किसी न किसी रूप में न्यूनतम हैं तथा दोनों पर किसी प्रकार का कोई प्रतिबंध (Constraints) नहीं होता है।

आरंभिक वृक्ष (आदि वृक्ष और गौण वृक्ष) कुछ निर्भरताओं (Dependencies) जैसे उपवर्गीकरण (Sub Categorization) निर्भरताओं और पूर्ति रिक्ति (Filler Gap) निर्भरताओं के लिए उपयुक्त व्यवहार क्षेत्र (Domains) हैं।

6.6.3 टैग में निहित एकीकरण और लक्षण-संरचना

वृक्ष-संलग्न व्याकरण (टैग) वस्तुतः लक्षण-संरचना पर आधारित रूपवाद है। इसमें संदर्भमुक्त व्याकरण की तुलना में स्थानिकता का अपेक्षाकृत बड़ा व्यवहार-क्षेत्र होता है। टैग के मूलभूत सिद्धांतों में इस बात की पूरी व्यवस्था है कि शब्दों में अंतर्निहित निर्भरताएँ (Dependencies) मूल वृक्षों में संपुटित (Encapsulated) हैं। किसी भी क्रिया और उसके कोणांक (Arguments) स्थानिक रूप में मूल वृक्ष में ही निहित होते हैं। जैसे सकर्मक क्रिया 'खाना' कर्म के रूप में किसी भोज्य पदार्थ की संज्ञापद (NP) के रूप में अपेक्षा करती है और 'सोना' अकर्मक क्रिया कर्म के रूप में किसी संज्ञापद की अपेक्षा नहीं करती। यह तथ्य उक्त क्रियाओं के मूलवृक्षों में ही निहित होता है।

टैग को आधारित (Embedded) करने का प्रमुख उद्देश्य इसी स्थानिकता या निर्भरताओं को पकड़ना है। इसके दो कारण हैं। एक, भाषावैज्ञानिक रुचि का विषय है कि स्थानिकता का क्षेत्र कितना व्यापक हो सकता है और दूसरा कारण यह है कि प्रतिबंधित लक्षणों के संचरण के कारण यह अभिकलनात्मक दृष्टि से (Computationally) भी अत्यंत कुशल रूपवाद (Formalism) है। इसीलिए लक्षण-संरचनाओं को मूल वृक्षों के साथ ही संबद्ध कर दिया जाता है। इसके विपरीत संदर्भमुक्त व्याकरण (CFG) में इन वृक्षों की शाखाओं को तोड़कर अलग-अलग नियमों में विभाजित कर दिया जाता है।

मूल वृक्षों के साथ संबद्ध लक्षण प्रणाली में प्रतिबंधों को निर्भर नोड के बीच सीधे वर्णित किया जा सकता है, इसलिए सामान्य वाक्य वाले आदि वाक्य में मुख्य क्रिया और संज्ञा-पद (क्योंकि दोनों ही उसी आदि वृक्ष के भाग हैं) की अन्विति समान हो सकती है। इसी प्रकार क्रिया के लिंग, वचन और पुरुष आदि को कर्ता के संज्ञापद के साथ अन्विति करने के लिए भी कुछ प्रतिबंधों को लागू किया जा सकता है। यही स्थिति उपवर्गीकरण (Sub-categorisation) के मामले पर भी लागू होती है। भाषा-वैज्ञानिक सिद्धांतों के अनुसार प्रत्येक क्रिया में उसके कोणांक (Arguments) अंतर्निहित होते हैं। जैसे 'खाना' क्रिया कर्म के रूप में एक संज्ञापद की अपेक्षा करती है।

वाक्य (1) - राम फल खाता है।

इसी प्रकार "देना" क्रिया दो कर्मों की अपेक्षा करती है।

प्रत्यक्ष कर्म और परोक्ष कर्म।

वाक्य (2) - राम श्याम को फल देता है।

टैग रूपवाद में अंतर्निहित तत्वों को संबंधित क्रिया के उपवर्गीकरण के रूप में मूलवृक्षों के भाग की तरह सन्निहित कर दिया जाता है।

अभिकलनात्मक भाषिकी में प्रयुक्त अधिकांश व्याकरण संबंधी रूपवाद (Grammar Formalisms) एकीकरण पर आधारित दृष्टिकोण अपनाते हैं क्योंकि लक्षण-संरचना का एकीकरण, संयोजन-लक्षण-संरचना (Combining Feature Structure) में प्रयुक्त

प्राथमिक संक्रिया (Primary Operation) है। एकीकरण का यह प्रत्यय वस्तुतः तर्कशास्त्र से लिया गया है। एकीकरण मुक्त क्रम (Order Free) की व्यवस्था है। इसके अंतर्गत भाषावैज्ञानिक को किसी भी अवयव पर कहीं भी प्रतिबंध लगाने की पूरी छूट होती है। विभिन्न रूपवादों (Formalisms) में लक्षण-संरचना का प्रत्यय भाषिक लक्षणों को वर्णित करने के लिए किया गया है। लक्षण-संरचना अनिवार्यतः लक्षण-मूल्य (Attribute-Value) के युग्म (Pairs) हैं। ये मूल्य आण्विक-प्रतीक (Atomic Symbols) या एक और लक्षण-संरचना भी हो सकते हैं। प्राकृतिक भाषा संसाधन (NLP) के अंतर्गत समीकरणों (Equations) का प्रयोग संदर्भयुक्त व्याकरण (CFG) के उपादानों के साथ किया जाता रहा है।

एकीकरण पर आधारित अधिकांश व्याकरण-रूपवादों को एकीकरण दृष्टिकोण पर आधारित संदर्भमुक्त व्याकरण का प्रत्ययमूलक चर (Notional Variation) माना जा सकता है। ये प्रणालियाँ बहुत सशक्त हैं और प्राकृतिक भाषा विन्यास के अधिकांश सिद्धांतों को व्यक्त करने में सक्षम हैं। सूचनापरक (Declarative) होने के कारण उनमें भाषिक गुणों को सीधे प्रकट करने की क्षमता है। एकीकरण के समीकरण में समानीकरण (सामान्यीकरण Generality) का गुण बहुत आकर्षक लगता है, लेकिन अभिकलनात्मक दृष्टि से (Computationally) यह बहुत कमजोर सिद्ध होता है। इस कमजोरी या अक्षमता का मुख्य कारण संदर्भमुक्त व्याकरण में 'स्थानिकता के व्यवहार क्षेत्र' (Domains of Locality) की सीमाएँ हैं। संदर्भमुक्त व्याकरण में एकीकरण प्रणाली के समीकरण अलग-अलग पुनर्लेखन नियमों के साथ संबद्ध हैं। इसलिए विभिन्न अवयवों के बीच के प्रतिबंधों की जाँच के लिए लक्षण-संरचना को व्यापक रूप में संचरित करना पड़ता है।

6.6.4 टैग की स्थानिकता का व्यवहार-क्षेत्र

अंतर्निहित भाषिक सिद्धांतों में यह परिकल्पना की गई है कि मूलवृक्षों में न्यूनतम भाषिक संरचना निहित होती है और इन निर्भरताओं को स्थानिक रूप में पकड़ा जा सकता है। यह भी परिकल्पना की गई है कि क्रिया और उसके कोणांक उसी मूलवृक्ष के भाग हैं। उसकी निर्भर मदों (Dependent Items) को आदिवृक्ष में या फिर गौणवृक्ष में, जिसके साथ उन्हें संलग्न किया जाता है, ढूँढ़ा जा सकता है। यह तभी होता है जब ऐच्छिक तत्वों (Adjunctions) को वृक्ष में प्रतिबंधों के तौर पर संनिहित किया जाता है। समीकरणों को मूलवृक्षों के नोडों की लक्षण-संरचना के रूप में वर्णित किया जाता है। उदाहरण के लिए, कर्ता और मुख्य क्रिया के अन्विति लक्षणों की पहचान मूलवृक्षों से हो जाती है। इस पहचान की जाँच शब्दों के स्तर पर ही की जा सकती है। इसके लिए लक्षण-संचरण (Feature Passing) की आवश्यकता नहीं होती।

6.6.5 टैग के अंतर्गत संलग्नक प्रतिबंध (Adjoining Constraints)

टैग के अंतर्गत विभिन्न प्रकार के प्रतिबंध लगाने की क्षमता है। इन प्रतिबंधों को मुख्यतः तीन भागों में विभाजित किया जा सकता है।

(क) अनिवार्य संलग्नक प्रतिबंध (Obligatory Adjoining Constraints) : ऊपरी और निचली संरचनाओं में लक्षणों के स्तर पर एकीकरण होने पर ही संलग्नक की संक्रिया संपन्न हो सकती है। यदि किसी गौणवृक्ष में कुछ ऐसे तत्व हों, जिनके बिना वाक्य संरचना बन

ही नहीं सकती, इसे अनिवार्य संलग्नक प्रतिबंध कहा जाता है। उदाहरण के लिए अंग्रेजी भाषा में Past Participle का प्रयोग has, had, have के बिना संभव नहीं है। निम्नलिखित वाक्य उक्त गौण क्रिया के बिना असंगत वाक्य कहलाएगा।

वाक्य - Ram gone to the school.

किंतु यदि इस वाक्य में 'has' का प्रयोग कर दिया जाए तो यह वाक्य व्याकरणिक दृष्टि से संगत या स्वीकार्य माना जाएगा।

वाक्य - Ram has gone to the school.

इस तत्त्व को यदि मूलवृक्षों के साथ ही लक्षण के रूप में प्रदर्शित कर दिया जाए तो Past Participle की संरचना को has आदि के बिना प्रजनित ही नहीं किया जा सकेगा।

(ख) चयनात्मक संलग्नक प्रतिबंध (Selective Adjoining Constraints) : कुछ संज्ञा पद ऐसे हैं, जिनके साथ विशेषण के रूप में किन्हीं खास शब्दों का ही प्रयोग किया जाता है। यद्यपि गलत शब्द रखने से अर्थ का अनर्थ तो नहीं होता लेकिन अटपटा-सा लगता है। उदाहरण के लिए, अंग्रेजी भाषा में handsome और beautiful विशेषणों का वितरण इस प्रकार किया जाता है कि handsome का प्रयोग सामान्यतः man के साथ और beautiful का प्रयोग woman के साथ किया जाता है। यदि हम चाहते हैं कि इस प्रयोग को नियमित बनाया जाए तो आवश्यक है कि इसकी व्यवस्था भी इनके मूलवृक्ष पर ही की जाए। टैग में इसके लिए चयनात्मक संलग्नक प्रतिबंध की व्यवस्था है।

(ग) शून्य संलग्नक प्रतिबंध (Null Adjoining Constraints) : यह प्रतिबंध टैग के अंतर्गत एक ऐसी व्यवस्था है जिससे अनेक गंभीर भाषिक अशुद्धियों से बचा जा सकता है। उदाहरण के लिए अंग्रेजी का यह पदबंध लिया जा सकता है। The tall good boy. शून्य प्रतिबंध की व्यवस्था न हो तो *The tall the good boy. पदबंध का प्रजनन भी किया जा सकता है। यह संरचना किसी भी रूप में स्वीकार्य नहीं हो सकती इसलिए आवश्यक है कि दूसरे विशेषण वाले N पर शून्य (NULL) प्रतिबंध लगाया जाए।

6.6.6 टैग और हिंदी व्याकरण

अंग्रेजी भाषा की तुलना में हिंदी भाषा को शब्दक्रम की दृष्टि से अपेक्षाकृत अधिक मुक्त क्रम की भाषा माना जा सकता है। उदाहरण के लिए वाक्य (1) देखें :

वाक्य (1) Ram killed Ravan.

अंग्रेजी के इस वाक्य में यदि शब्दों का क्रम बदल दिया जाए तो अर्थ ही बदल जाता है। [देखें, वाक्य (2)]

वाक्य (2) Ravan killed Ram.

इसके विपरीत हिंदी, संस्कृत और अन्य भारतीय भाषाओं में शब्दों का क्रम बदलने पर भी अर्थ में परिवर्तन नहीं होता,। देखें वाक्य (3), (4), (5) और (6)।

वाक्य (3) राम ने रावण को मारा।

वाक्य (4) रावण को राम ने मारा।

वाक्य (5) मारा रावण को राम ने।

वाक्य (6) मारा राम ने रावण को।

संगतः इसका मुख्य कारण कादचित्त यही है कि विभक्ति चिह्नों के व्यापक प्रयोग के कारण शब्दों के क्रम में परिवर्तन के बावजूद वाक्य के संपूर्ण अर्थ में विशेष परिवर्तन नहीं होता, जबकि अंग्रेजी में अर्थ का अनर्थ होने की आशंका रहती है।

यदि विभक्ति चिह्नों में परिवर्तन कर दिया जाए तो निश्चय ही वाक्य का अर्थ बदल जाएगा। देखें, वाक्य (7)

वाक्य (7) रावण ने राम को मारा।

वस्तुतः यह कारक चिह्न उपर्युक्त वाक्य का नाभिकीय या अनिवार्य तत्व है। इसके प्रयोग के बिना यह वाक्य अधूरा है। किंतु अनेक स्थलों पर वक्ता के मन में मौजूद होते हुए भी इनका प्रयोग नहीं किया जाता। लीच (1969) इसे शून्य अभिव्यक्ति (Zero Expression) का नाम देते हैं। देखें निम्नलिखित वाक्य :

वाक्य - राम घर पहुँचा।

वाक्य - Ram reached home.

डॉ. सूरजभान सिंह (1985) के अनुसार, 'जाना', 'आना', 'पहुँचना' और 'लौटना' आदि गंतव्य सूचक क्रियाओं से पूर्व यदि गंतव्य स्थल का बोध समग्रता में व्यक्त हो तो सामान्यतः 'में', 'पर' आदि परसर्गों का प्रयोग नहीं किया जाता। यदि गंतव्य स्थल का बोध किसी स्थल विशेष के अंतरंग में अवस्थिति या प्रविष्टि के रूप में व्यक्त हो तो परसर्ग 'में', 'पर' आदि का प्रयोग होता है। जैसे,

वाक्य - वह कमरे में पहुँचा।

वाक्य - He reached in the room.

सभी भाषाओं में कुछ प्रवृत्तियाँ समान रूप से अनिवार्य और ऐच्छिक घटक के रूप में क्रियापद में ही अंतर्निहित होती हैं, किंतु कुछ प्रवृत्तियाँ ऐसी भी होती हैं जो भाषा विशेष में ही परिलक्षित होती हैं।

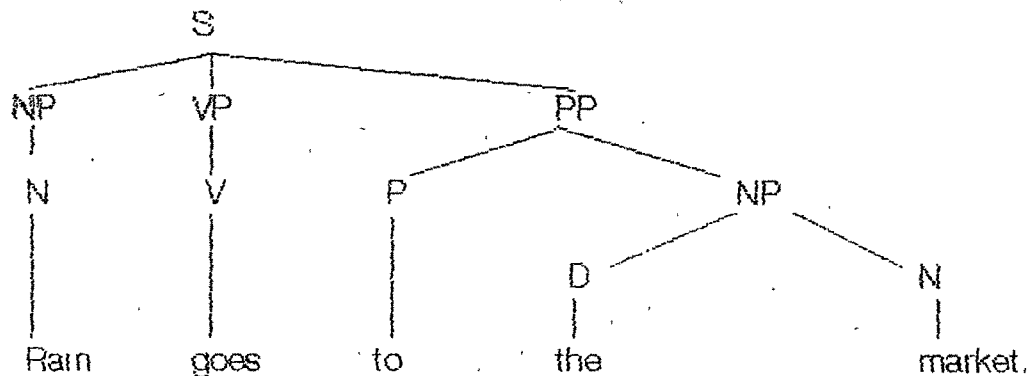
हिंदी और अन्य भारतीय भाषाओं में परसर्गों के प्रयोग की बहुलता के कारण शब्दों का क्रम बदलने पर भी वाक्यार्थ में विशेष परिवर्तन नहीं होता, किंतु हिंदी में परसर्गों का आधार बहुत स्पष्ट नहीं है।

सर्वभाषा व्याकरण (Universal Grammar) के अंतर्गत प्रयुक्त परसर्ग सार्वभौमिक कहे जा सकते हैं और हिंदी में भी इस प्रकार के अनेक परसर्गों का प्रयोग किया जाता है। जैसे

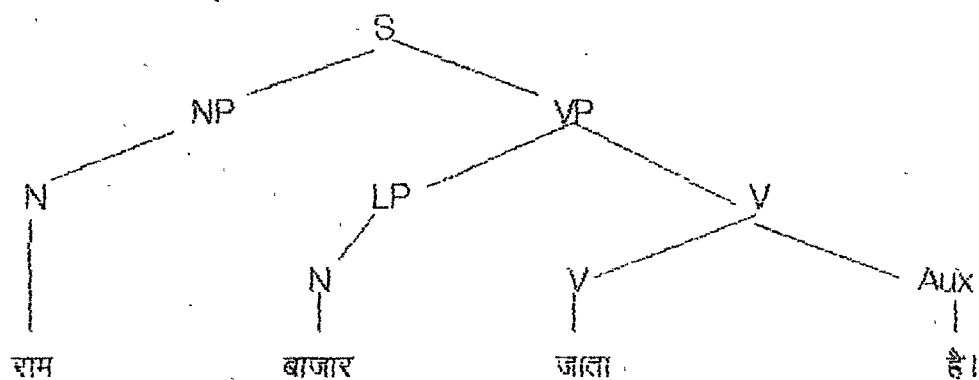
करण, अपादान, अधिकरण आदि। ऐसी स्थिति में हिंदी और अंग्रेजी में व्याकरण के स्तर पर पर्याप्त भिन्नता होने के कारण वृक्ष संलग्न व्याकरण (TAG) के माध्यम से पदनिरूपण (Parsing) तब तक संभव नहीं होगा जब तक कि दोनों भाषाओं के गैर टर्मिनल प्रतीकों को संबंधित भाषाओं के अनुरूप निर्धारित नहीं कर लिया जाता। अंग्रेजी भाषा में टैग के अंतर्गत पदनिरूपण के लिए तीन प्रमुख पदबंधों की परिकल्पना की गई है। NP (संज्ञा पदबंध), VP (क्रिया पदबंध) और PP (परसर्ग पदबंध)। उदाहरण के रूप में निम्नलिखित वाक्य देखें :

वाक्य - Ram goes to the market.

टैग व्याकरण के अंतर्गत इसकी संरचना इस प्रकार होगी :



किंतु डॉ हेमंत दरबारी (1991) के अनुसार हिंदी, संस्कृत या अन्य भारतीय भाषाओं में इसका पदनिरूपण इस प्रकार होगा।



बीज वाक्यों का कोशीय अंतरण :

सीमित कॉर्पस के आधार पर अंग्रेजी से हिंदी और हिंदी से अंग्रेजी में कोशीय अंतरण के लिए टैग को अपनाने का मुख्य कारण यह है कि अंग्रेजी और हिंदी इन दोनों भाषाओं के वाक्य विन्यास में पर्याप्त भिन्नता है। अंग्रेजी भाषा के वाक्य-विन्यास में शब्दों का क्रम अपेक्षाकृत निश्चित है और उसमें परिवर्तन करने से अर्थ का अनर्थ भी हो सकता है। इसके विपरीत हिंदी अपेक्षाकृत मुक्त शब्द क्रम की भाषा है। इसमें शब्दों का कर्म बदलने से अर्थ का अनर्थ नहीं होता और अर्थ में भी कोई विशेष परिवर्तन नहीं होता।

भारतीय भाषाओं में अंतर्निहित समानता है। कदाचित् इसी समानता के कारण भारतीय भाषाओं में कोशीय अंतरण की संभावनाएं अपेक्षाकृत अधिक है। आई.आई.टी., कानपुर द्वारा विकसित 'अनुसारक' के माध्यम से भारतीय भाषाओं में परस्पर अनुवाद की

संभावनाएं उजागर हुई हैं, किंतु दो विभिन्न वाक्य-विन्यास की भाषाओं के बीच अनुवाद का कोई सफल रूपवाद अभी तक पूर्णतः विकसित नहीं हो पाया है। प्राकृतिक भाषाओं का व्याकरण गणित के समीकरणों के समान इतना व्यवस्थित नहीं होता, जितना कि ऊपर से दिखाई पड़ता है। उदाहरण के लिए, हिंदी में परसर्गों की व्यवस्था इतनी जटिल और भ्रामक है कि इससे समस्या सुलझाने के बजाय उलझाने के अधिक आसार दिखाई पड़ते हैं। उदाहरण के लिए हिंदी में 'को' और 'से' परसर्गों का प्रयोग इतना यादृच्छिक है कि उसके लिए कारक व्यवस्था का कोई नियमित आधार ऊपर से दिखाई नहीं पड़ता। देखें

वाक्य	राम श्याम को पीटता है।
वाक्य	राम को बुखार है।
वाक्य	पेड़ से पत्ता गिरता है।
वाक्य	राम से चला नहीं जाता।
वाक्य	वह चाकू से फल काटता है।
वाक्य	राम श्याम से मिलता है।

इन उदाहरणों से स्पष्ट है कि जब तक हिंदी में परसर्गों की व्यवस्था का स्पष्ट और नियमित आधार नहीं मिलता तब तक इन वाक्यों का पदनिरूपण नहीं किया जा सकता और जब तक पदनिरूपण नहीं होता, तब तक अँग्रेजी-हिंदी में कोशीय अंतरण नहीं हो सकता।

प्राकृतिक भाषाओं में प्रयुक्त सभी प्रकार के वाक्यों का मूल आधार बीज वाक्य ही होते हैं। यदि इन बीज वाक्यों में प्रयुक्त पदबंधों की व्याकरणिक कोटियों की पहचान कर ली जाए और संयोजन, लोप, रूपांतरण तथा आगम प्रक्रियाओं को रेखांकित किया जा सके तो प्राकृतिक भाषाओं के सभी वाक्यों का पदनिरूपण किया जा सकता है। प्रो. सूरजभान सिंह (1985) के अनुसार हिंदी भाषा में प्रयुक्त सभी सार्वभौमिक तथा भाषाविशिष्ट तत्वों को आधार बनाकर उन्हें 14 बीज वाक्यों में समाहित किया जा सकता है।

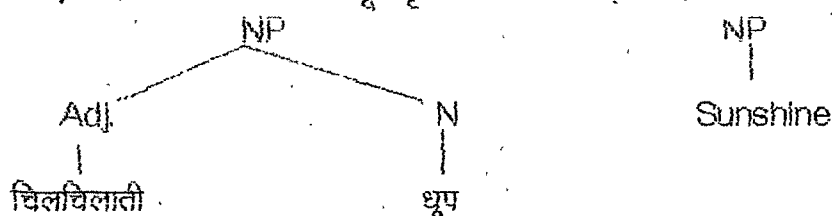
अंतरण शब्दवृत्त (Transfer Lexicon) :

इस रूपवाद के माध्यम से कोशीय अंतरण के लिए न तो नियम से नियम का अंतरण होता है और न ही शब्द से शब्द का। इसके अंतर्गत अंतरण होता है शब्दवृक्ष से शब्दवृक्ष का। उदाहरण के लिए हिंदी और अँग्रेजी के दो वाक्यों को लिया जा सकता है :

वाक्य	राम को बुखार है।
वाक्य	Ram has a fever.

दोनों के अंतर्गत कोशीय अंतरण के लिए इन शब्दवृक्षों को शब्दवृत्त (Lexicon) के अंतर्गत एक साथ रखा जाता है। इसलिए पर्याप्त भिन्नता के बावजूद व्युत्पन्न वृक्ष के रूप में दोनों भाषाओं का पद निरूपण सरलता से हो जाता है और फिर प्रजनन (Generation) में भी कोई कठिनाई नहीं होती। दो भाषाओं के मूलवृक्षों को एक साथ जिस स्थान पर रखा जाता है उसे 'अंतरण शब्दवृत्त' (Lexicon Transfer) कहा जाता है।

इसी प्रकार अभिव्यक्तियों, वाक्यों, मुहावरों आदि के शब्दगुच्छों को भी शब्दवृक्ष के रूप में एक साथ रखा जा सकता है। टैग की भाषा में इन्हें Multi Component Anchor अर्थात् बहु-उपादानीय तंजर कहा जाता है। वस्तुतः सहप्रयोग और मुहावरों में इनका खास तौर पर प्रयोग किया जाता है। उदाहरण के लिए 'चिलचिलाती धूप' में 'चिलचिलाती' शब्द धूप के सहप्रयोग के रूप में आया है। निश्चय ही यह विशेषण है, इसका प्रयोग 'धूप' के अलावा किसी अन्य संज्ञापद के साथ नहीं किया जाता। इसलिए आवश्यक है कि धूप के अनिवार्य विशेषण के रूप में इसे मूलवृक्ष में ही प्रदर्शित किया जाए। इसके विपरीत अंग्रेजी में इसके लिए अनेक विशेषणों का प्रयोग किया जा सकता है। इसलिए सहप्रयोग के रूप में मूलवृक्ष के अंतर्गत इसे दिखाने की आवश्यकता नहीं है।



अंतरण शब्दवृक्ष के अंतर्गत निम्नलिखित का मिलान किया जाता है :

- शब्द परिवारों का मिलान।
- शब्द वृक्षों का मिलान।
- नोडों का मिलान और नोडों के माध्यम से नोडों के साथ संबद्ध लक्षणों का मिलान।¹⁴

6.7 अध्याय 6 की संदर्भ सूची

1. कंप्यूटर के भाषिक अनुप्रयोग, विजय कुमार मल्होत्रा, पृ. 183-190
2. कंप्यूटर के भाषिक अनुप्रयोग, विजय कुमार मल्होत्रा, पृ. 15-24
3. ड्रीमलैंड का इलस्ट्रेटेड कंप्यूटर एनसाइक्लोपीडिया, प्रशांत गुप्ता, पृ.601
4. ड्रीमलैंड का इलस्ट्रेटेड कंप्यूटर एनसाइक्लोपीडिया, प्रशांत गुप्ता, पृ.314
5. सी डेक का पैम्फलेट
6. CHIPहिंदी, अप्रैल 2000, पृ.13
7. सी-डेक का पैम्फलेट - मल्टी-लिंग्वल सॉफ्टवेयर डेवलपमेंट टूल्स
8. वीपीवी लीप ऑफिस, पृ. 252
9. आई.एस.एम.2000 मैनुअल, पृ.39
10. कंप्यूटर संचार सूचना, जून, 2003, पृ.24
11. आवाज पहचानक : कैसे और क्या, सुशांत कुमार, कंप्यूटर संचार सूचना, अक्टूबर 2001, पृ. 66-68
12. आवाज पहचानक, कंप्यूटर संचार सूचना, जून 2003, पृ. 54-56
13. कृत्रिम बुद्धि, के.डी.पावटे, पृ. 85-91
14. कंप्यूटर के भाषिक अनुप्रयोग, विजय मल्होत्रा, पृ. 205-245