

УДК 519.226:534.78:681.327

СИСТЕМА РАСПОЗНАВАНИЯ ИЗОЛИРОВАННЫХ СЛОВ ПО ЧАСТНОЙ
АВТОКОРРЕЛЯЦИОННОЙ ФУНКЦИИ

А.В.Кельманов

Целью данной работы является построение системы распознавания изолированных слов, в которой в качестве признаков используется частная автокорреляционная функция. Анализ речевого сигнала во временной области при помощи модели авторегрессии был успешно применен для эффективного сжатия описания речевого сигнала [1,2]. В этой связи появились попытки построения системы распознавания непосредственно по параметрам авторегрессии [3]. Однако эти попытки не дали ожидаемых результатов. Было показано [4], что высокую надежность распознавания по параметрам авторегрессии можно получить лишь в том случае, если будет учтена статистическая зависимость этих параметров.

Впервые частная автокорреляционная функция как система признаков была использована в [3] для построения системы распознавания изолированных слов, включающей 10 цифр. Опубликованные результаты подают определенную надежду на возможность использования этой системы признаков. Тем не менее остается открытым вопрос о ее применимости для системы распознавания с большим словарем не только японским, английским, но и русским, а также вопрос о выборе подходящих моделей авторегрессии. Данную работу следует рассматривать как ответ на эти вопросы.

§1. Система признаков

Пусть речевой сигнал $s(t)$ проквантован с частотой квантования $F = 1/T$, так что $s_d = s(nT)$, где T — интервал квантования. Обозначим через K объем словаря системы распознавания изолирован-

ных слов. Каждое k -е слово словаря разобьем на L_k сегментов $\{a_n\}$, $n = \overline{1, N}$. Будем описывать каждый 1-й сегмент $\{a_n\}^{(1)}$ моделью авторегрессии p -го порядка, имеющей вид:

$$\Lambda^{(1)}(z^{-1})w_n^{(1)} = \alpha_n^{(1)}, \quad w_n^{(1)} = \nabla_n^{(1)}a_n^{(1)},$$

$$\Lambda^{(1)}(z^{-1}) = 1 - \sum_{i=1}^p a_i^{(1)} z^{-i}$$

- стационарный оператор авторегрессии p -го порядка, $\nabla_n^{(1)} = 1 - a^{(1)}z^{-1}$ - адаптивный оператор взятия разности первого порядка, $a^{(1)} = r_1^{(1)}(a)/r_0^{(1)}(a)$, $r_0^{(1)}(a)$, $r_1^{(1)}(a)$ - автокорреляции сигнала $\{a_n\}^{(1)}$ для нулевой и первой задержек, $\alpha_n^{(1)}$ - белый шум с дисперсией $[\sigma_\alpha^{(1)}]^2$.

Оценки параметров оператора $\Lambda^{(1)}(z^{-1})$ находятся из решения системы нормальных уравнений Ила-Уокера рекуррентным методом Дурбина [?]. Соотношения имеют следующий вид:

$$\hat{a}_{j+1,1}^{(1)} = \hat{a}_{j1}^{(1)} - \hat{a}_{j+1,j+1}^{(1)} \hat{a}_{j,j-1+1}^{(1)}, \quad j = 1, 2, \dots, J, \quad (I)$$

$$\hat{a}_{j+1,j+1}^{(1)} = \frac{r_{j+1}^{(1)} - \sum_{i=1}^j \hat{a}_{ji}^{(1)} \hat{r}_{j-1+1}^{(1)}}{1 - \sum_{i=1}^j \hat{a}_{ji}^{(1)} \hat{r}_i^{(1)}},$$

где первый индекс соответствует очередному шагу при оценивании, второй - номеру параметра оператора $\Lambda^{(1)}(z^{-1})$, а $\hat{r}_i^{(1)}$ - оценка автокорреляционной функции сигнала $\{w_n\}^{(1)}$, $n = \overline{1, N-1}$:

$$\hat{r}_i^{(1)} = \frac{\sum_{n=1}^{N-1-i} w_n^{(1)} w_{n+i}^{(1)}}{\sum_{n=1}^{N-1} [w_n^{(1)}]^2}.$$

Величина a_{11} , рассматриваемая как функция задержки i , называется частной автокорреляционной функцией.

Каждый 1-й сегмент слова будем характеризовать вектором параметров $\hat{a}_{pp}^{(1)} = (\hat{a}_{11}^{(1)}, \hat{a}_{22}^{(1)}, \dots, \hat{a}_{pp}^{(1)})$, а k -е слово словаря - по - следовательно векторов $\{a_{pp}^{(1)}\}$, $1 = \overline{1, L_k}$. Построенная система признаков обладает следующими свойствами.

1. Из (I) видно, что рекуррентная процедура решения позволяет последовательно оценивать параметры операторов авторегрессии $(j+1)$ для $j = 0, p-1$. При этом, используя теорию ортогональных многочленов [8], можно показать, что полиномы $A_{j+1}^{(1)}(z^{-1})$ комплексной переменной z^{-1} ортогональны на единичной окружности.

2. Для процесса авторегрессии p -го порядка $a_{11} = 0$ при $i > p$ [7].

3. $|a_{11}| < 1$ (см. [5]).

4. Параметры a_{11} статистически независимы [7].

5. Из (I) следует, что при увеличении порядка модели авторегрессии с i на $i + 1$ первые i компонент вектора $\hat{a}_{i+1, i+1}^{(1)}$ образуют вектор, равный вектору $a_{11}^{(1)}$, т.е. не изменяются с ростом порядка модели авторегрессии в отличие от параметров последней.

§2. Мера сходства и алгоритм распознавания

По аналогии с [9] введем меру сходства между сегментами 1 и q :

$$d_{1q} = \frac{\alpha^2}{\alpha^2 + \rho_{1q}^2},$$

где $\rho_{1q}^2 = \sum_{i=1}^p [a_{11}^{(1)} - a_{11}^{(q)}]^2$ - квадрат евклидова расстояния между векторами $a_{pp}^{(1)}$ и $a_{pp}^{(q)}$, α - экспериментально подбираемая константа.

Членение слова на сегменты проводится при помощи метода разбиения с перекрытием при длительности сегмента $T_a = T(N-1)$ и сдвиге от сегмента к сегменту на время $\Delta T = T(\Delta N - 1)$. Так что число сегментов в слове находится из соотношения:

$$L_{сл} = \left\lfloor \frac{T_{сл} - T_a}{\Delta T} \right\rfloor + 1 = \left\lfloor \frac{N_{сл} - N}{\Delta N} \right\rfloor + 1,$$

где N - число отсчетов в слове, T - длительность слова. При этом остаток слова отбрасывается, если его длительность меньше чем T_a .

Вычисление меры сходства D_{max} между контрольным и эталонным словами производится по методу динамического программирования, описанному в [9]. Окончательная мера сходства D между двумя словами

k и m , содержащими L_k и L_m сегментов соответственно, определяется путем нормирования по длине слова - делением $D_{\max}$ на длину более длинного слова: $D = D_{\max}/L_{\max}$, $L_{\max} = \max(L_k, L_m)$.

Каждое m -е контрольное слово сравнивается со всеми эталонными, так что в результате формируется последовательность мер сходства $D_1^{(m)}, \dots, D_K^{(m)}$. Окончательное решение принимается по правилу:

$$k = \underset{i=1,2,\dots,K}{\operatorname{argmax}} D_i^{(m)}. \quad (2)$$

Результат распознавания считается верным, если $k=m$ (предполагается, что m -е контрольное слово соответствует k -му эталону). Распознавание одного слова длины L из словаря, содержащего K слов, требует выполнения

$$\tau = c K L L_{\text{cp}}^{(3)} \quad (3)$$

операций, где $c = \operatorname{const}$,

$$L_{\text{cp}}^{(3)} = \frac{1}{K} \sum_{k=1}^K L_k^{(3)}$$

- средняя длина эталонного слова, а $L_k^{(3)}$ - длина k -го эталонного слова. Из (3) видно, что трудоемкость алгоритма пропорциональна квадрату длины слова и линейно возрастает с увеличением порядка модели и объема словаря. Следовательно, распознавание больших словарей потребует большого расхода машинного времени. Поэтому необходимо заботиться о его сокращении.

§4. Результаты эксперимента

Речевой материал. При проведении эксперимента использовался словарь, состоящий из 100 русских слов (табл. I), разработанный В.М.Белявским и рекомендованный Всесоюзной школой-семинаром "Автоматическое распознавание слуховых образов" для проверки системы распознавания изолированных слов. По данному словарю одним диктором-мужчиной были сформированы две последовательности - эталонная и контрольная - с интервалом произнесения полгода.

Условия эксперимента. Ввод изолированных слов в ЭВМ "Минск-32" осуществлялся непосредственно в машинном зале (уровень шума ~ 60 дБ) через микрофон типа МД-59 и семипразрядный аналого-цифровой преобразователь с частотой квантования 20 кГц. После ввода производилось автоматическое определение границ слов и запись введенных слов на магнитную ленту.

Т а б л и ц а I

1. Кефир	26. Борщ	51. Плохой	76. Монета
2. Инфический	27. Гиря	52. Похвала	77. Танец
3. Графин	28. Богема	53. Уходить	78. Кармен
4. Сохи	29. Ген	54. Повемка	79. Мы
5. Химик	30. Гиды	55. Бузина	80. Домой
6. Ляхи	31. Сами	56. Раздетый	81. Если
7. Михей	32. Домино	57. Казенный	82. Разъезд
8. Кишки	33. Онеметь	58. Резец	83. Головка
9. Лакей	34. Извинись	59. Зеркало	84. Засовы
10. Чуть	35. Обновить	60. Петь	85. Зову
11. Щек	36. Навес	61. Пилон	86. В этом
12. Няня	37. Овин	62. Отпилить	87. Свободный
13. Кляча	38. Советский	63. Отбилсь	88. Во рву
14. Ладья	39. Каждый	64. Регби	89. Симфазно
15. Грязь	40. Подожди	65. Себе	90. Борьба
16. Ручной	41. Пожар	66. Ребята	91. Бурить
17. Ночевать	42. Зажим	67. Шабаш	92. Совсем
18. Мальчик	43. Железо	68. Галоши	93. Посуда
19. Бычок	44. Пижон	69. Шило	94. Один
20. Куча	45. Уже	70. Леший	95. Два
21. Гуца	46. Прижаться	71. Прошу	96. Четыре
22. Тощий	47. Охалка	72. Шуба	97. Газон
23. Ищите	48. Орех	73. Левша	98. Фуга
24. Барщина	49. Доха	74. Коневод	99. Где
25. Мужчина	50. Хам	75. Тунец	100. Дереза

У с л о в и я а н а л и з а. Акустические реализации слов, хранящиеся на магнитной ленте, обрабатывались с целью выделения признаков при помощи процедуры, описанной выше. Анализ проводился по методу разбегания с перекрытием при длительности сегмента $T_a = 25,6$ мсек и сдвиге от сегмента к сегменту на $\Delta T = 16$ мсек. Каждый сегмент взвешивался окном Хэмминга и описывался моделью авторегрессии 30-го порядка, а на магнитную ленту были записаны последовательности векторов $\{a_{30,30}\}$.

Р е з у л ь т а т ы р а с п о з н а в а н и я. Процедура распознавания, описанная выше, была опробована для моделей авторегрессии с различными значениями порядка p . Результаты распознавания приведены в табл. I и на рис. I.

Таблица 2

Порядок модели	Число правильно распознанных слов	Надежность, %
2	34	34
4	61	61
6	79	79
8	88	88
10	93	93
12	97	97
14	99	99
16	100	100
18	100	100
20	100	100
22	100	100
24	100	100
26	100	100
28	100	100
30	100	100

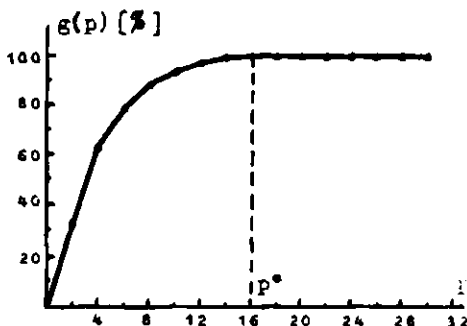


Рис. 1. Зависимость надежности распознавания от порядка модели.

Из рис.1 видно, что надежность распознавания $g(p)$ является монотонно возрастающей функцией порядка модели p . Этот экспериментальный результат хорошо согласуется с известным [7] теоретическим, состоящим в том, что среднеквадратичная

ошибка аппроксимации сигнала моделью авторегрессии является монотонно убывающей функцией порядка модели p .

Значение константы α было выбрано равным 1.

Сокращение трудоемкости. Известно несколько эвристических методов сокращения времени принятия решения [4,6,9,10]. В данной работе предлагается еще один, который в сочетании с указанным позволит дополнительно сократить время счета. В его основе лежит свойство монотонности функции $g(p)$ и свойство 5 - частной автокорреляционной функции.

Обозначим через p^* такое значение порядка модели p , что $g(p) = \text{const}$ при $p \geq p^*$ и $g(p) < \text{const}$ при $p < p^*$. В данном случае (см.рис.1) $p^* = 16$, $\text{const} = 100\%$.

Пусть при распознавании m -го контрольного слова получена последовательность мер сходства $\{D_1^{(m)}\}$. Упорядочим ее по убыванию. Очевидно, что после сортировки при $p \geq p^*$ в соответствии с (2) будет получена последовательность $D_1^{(m)}, \dots, D_p^{(m)}$, в которой на первом месте стоит мера сходства для m -го контрольного слова. При $p = p^*, p < p^*$ получим некоторую последовательность $D_1^{(m)}, \dots, D_p^{(m)}, \dots, D_{p^*}^{(m)}$, в которой мера сходства m -го контрольного

слова с m -м эталонным стоит на некотором K_1 -м месте. Следовательно, возможно сокращение времени распознавания за счет уменьшения числа претендентов.

Безусловно, вывод о возможности сокращения времени счета зависит от вида функции $K_1(p)$, которая нам неизвестна.

В соответствии с описанной процедурой трудоемкость алгоритма (2) можно переписать в виде:

$$\tau' = cK_{LL}^{(a)} p_1 + cK_{LL}^{(a)} p^*,$$

где $\bar{L}_{cr}^{(a)}$ — средняя длина эталонного слова по первым K_1 претендентам, вычисляемая по аналогии с (4). Предполагая, что $\bar{L}_{cr}^{(a)} = \bar{L}_{cr}^{(a)}$,

(5) можно переписать в виде $\tau' = c\bar{L}_{cr}^{(a)} (Kp_1 + K_1 p^*)$.

Таким образом, задачу сокращения трудоемкости алгоритма можно сформулировать в виде задачи целочисленного программирования:

$$(K_1, p_1) = \underset{\substack{0 < p_1 \leq p^* \\ 0 < K_1 \leq K_1(p)}}{\operatorname{argmin}} \tau' = \underset{\substack{0 < p_1 \leq p^* \\ 0 < K_1 \leq K_1(p)}}{\operatorname{argmin}} (Kp_1 + K_1 p^*) \quad K_1 = K_1(p)$$

Т а б л и ц а 3

Порядок модели p	$K_1(p)$
2	64
4	23
6	10
8	9
10	4
12	3
14	3
16	1
18	1
20	1
22	1
24	1
26	1
28	1
30	1

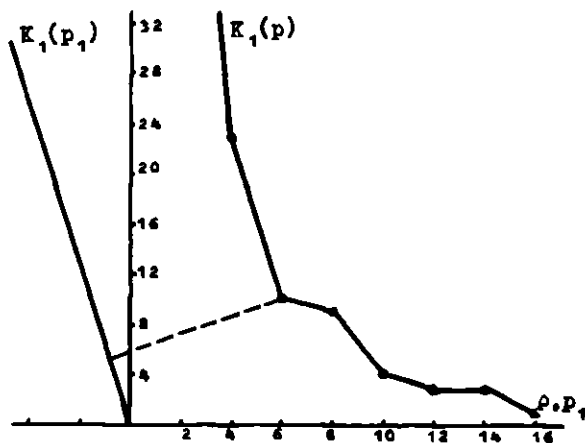


Рис.2. График функции $K_1(p)$ и $K_1(p_1)$.

Параллельно с экспериментом по распознаванию были получены значения функций $K_1(p)$, приведенные в табл. 3 и на рис. 2.

На рис. 2 также приведен график прямой, задаваемой уравнением $Kr_1 + K_1p^* = 0$ (при $K = 100$, $p^* = 16$). Что касается времени распознавания, то оно пропорционально расстоянию между прямой $K_1(p_1)$ и ломаной $K_1(p)$ (с учетом целочисленных ограничений). Из рис. 2 видно, что $\tau' = \tau_{\min}$ при $p_1 = 6$ и $K_1 = 10$. В таком случае время принятия решения сокращается более чем в 2 раза. В соответствии с описанной процедурой сокращения времени принятия решения эксперимент по распознаванию был повторен для $p_1 = 6$ и $p^* = 16$. При этом надежность распознавания осталась прежней (100%).

§4. Обсуждение результатов

Основной целью данной работы было исследование возможности использования частных автокорреляций в качестве признаков и меры сходства, предложенной в [9], для построения системы распознавания. Поэтому в работе умышленно не были задействованы некоторые из параметров, применявшиеся в других системах распознавания изолированных слов и хорошо зарекомендовавшие себя. В табл. 4 приведены данные построенной системы, а также известных зарубежных аналогов.

В работах [4, 6] в качестве признаков были использованы непосредственно параметры авторегрессии, а число признаков было на единицу больше порядка модели и составило величину, равную 7 в [4] и 15 в [6]. В качестве меры сходства служила мера, введенная в [4], учитывающая статистическую зависимость параметров авторегрессии и отличная от меры, применявшейся в [3] и в данной работе.

В системе распознавания, исследованной в [3], использовалась мера, построенная на основе евклидовой метрики и для параметров авторегрессии, и для частной автокорреляции. Из табл. 4 видно, что в этом случае непосредственное применение параметров авторегрессии дает надежность существенно меньшую, чем при использовании частной автокорреляции даже при том условии, что размерность пространства в первом случае выше. Заметим, что автором данной работы ранее также была предпринята аналогичная попытка построения системы по параметрам авторегрессии для $p = 12$ и получена надежность распознавания 33%. Для сравнения отметим, что частная автокорреляционная функция при том же p позволяет получить надежность распознавания 97% (табл. 2). Эти факты говорят о том, что при построении

Т а б л и ц а 4

Система	Словарь/объем	Признаки	Р _{кв} [кгц]	Шум [дБ]	Микрофон	Число разрядов преобразователя	Т _а [мсек]	ΔТ [мсек]	Надежность
Ф.Ита-кура [3]	Алфавитно-цифровой/36 Японские географические названия/200	Параметры модели АР (6)	6,667	68	телефон	-	30	15	88,6 98,95
Д.Уайт [5]	Алфавитно-цифровой/36 Североамериканские штаты/91	Параметры модели АР (14)	10	65	МБГ	8	25,6	12,8	98,0 99,6
А.Иичикава [2]	Цифры/10 Цифры/10	Параметры модели А (7) Частная авторегрессия для модели АР (5)	10	-	-	II	25,6	10	77 98
.	Частотный/100	Частная авторегрессия для модели АР (16)	20	60	МД-59	7	25,6	16	100

Примечание: АР(6) - это авторегрессия 6-го порядка и т.д.

системы по параметрам авторегрессии нельзя использовать меру сходства, построенную на основе евклидовой метрики.

Кроме того, необходимо отметить, что в [3] зависимость надежности распознавания от порядка модели авторегрессии для частной автокорреляционной функции не носит монотонного характера, как это должно быть теоретически. Этот результат объясняется тем, что в [3] при оценивании параметров проводились вычисления с фиксированной запятой, что понизило точность вычислений и привело к ошибкам. В этой же работе приведен результат, полученный по сглаженным оценкам частной автокорреляции. После сглаживания надежность распознавания стала монотонной функцией, а ее значение при $p = 5$ повысилось с 98% до 100%. В данной работе, как видно из рис.1, надежность распознавания получается монотонной без сглаживания оценок, так как вычисления проводились с плавающей запятой.

Отличительной особенностью системы признаков, предлагаемой в данной работе, является то, что на этапе анализа сигнал описывается моделью авторегрессии со стационарной адаптивной разностью первого порядка в то время, как в [3] для этой цели служила обычная модель авторегрессии.

Сокращение трудоемкости описанным методом в общем случае может привести к некоторому понижению надежности распознавания по следующим причинам. Во-первых, функция $K_1(p)$ априори неизвестна, а строится после однократного эксперимента. Во-вторых, при повторном испытании для того же диктора эта функция может измениться таким образом, что распознаваемое слово при $p = p_1 \leq p^*$ не попадет в сокращенное число претендентов, что приведет к ошибке. Однако, учитывая гипотезу компактности, принятую во всех работах по распознаванию, можно считать, что для одного и того же диктора функция $K_1(p)$ при повторных экспериментах меняется незначительно. Наконец, можно построить $K_1(p)$ по результатам нескольких экспериментов. Метод построения очевиден.

Распознавание одного слова из словаря в 100 слов занимает на ЭВМ "Минск-32" в среднем около 7 мин с учетом работы внешних устройств (магнитные ленты, печать).

Перейдем теперь к вопросу о выборе порядка модели. Из (3) видно, что за счет сокращения размерности пространства, т.е. уменьшения порядка модели авторегрессии, можно добиться сокращения трудоемкости. Поэтому необходимо позаботиться об уменьшении p .

Рассмотрим сначала физические требования, предъявляемые к порядку модели при анализе сигнала. При частоте квантования 10 кГц,

когда полоса частот речевого сигнала ограничивается частотой 5 кГц, энергетический спектр обычно [11] содержит 5 максимумов, соответствующих 5 формантным колебаниям, каждое из которых описывается однократным резонатором, который, в свою очередь, представим в виде модели $AR(2)$. Таким образом, для описания пятиформантного сигнала необходима модель $AR(10)$. Для учета влияния голосового источника порядок необходимо повысить. Обычно "оптимальным" считается значение $p = 12-15$. При увеличении частоты квантования сигнала вдвое порядок модели также должен быть увеличен по крайней мере в 2 раза. В самом деле, если при частоте квантования 10 кГц каждый отсчет сигнала зависел только от p предыдущих, то при частоте квантования 20 кГц "расстояние" между отсчетами увеличивается вдвое и, таким образом, каждый отсчет будет зависеть от $2p$ предыдущих. Следовательно, при частоте квантования 20 кГц необходимо по крайней мере 24-30 параметров. Если к тому же учесть, что в этом случае спектр сигнала ограничивается частотой 10 кГц, а в диапазоне от 5 до 10 кГц есть дополнительные максимумы, то становится очевидным, что порядок модели необходимо дополнительно повысить.

Вернемся теперь к результатам эксперимента. Из него следует, что при построении системы при частоте квантования 20 кГц можно обойтись моделью $AR(16)$, т.е. довольно грубым анализом как в спектральной, так и во временной областях. В то же время нельзя не учитывать тот факт, что с увеличением объема словаря значение p^* может повыситься. Окончательные выводы о порядке модели можно будет сделать лишь после проведения эксперимента по фонемному перекодированию речи.

1. При построении систем распознавания изолированных слов по параметрам моделей авторегрессии необходимо учитывать, что статистическая независимость частных автокорреляций в отличие от параметров авторегрессии позволяет использовать меру сходства между сегментами речевого сигнала, построенную на основе евклидовой метрики.

2. Надежность распознавания по частной автокорреляции является монотонно возрастающей функцией порядка аппроксимирующей модели.

3. Для систем распознавания, аналогичных рассмотренной, порядок модели при частоте квантования 20 кГц должен быть не ниже 16. Разработанная система распознавания изолированных слов при использовании 16 частных автокорреляций и частоте квантования 20 кГц, позволяет работать одному диктору с высокой надежностью.

4. Положительные результаты экспериментов свидетельствует о том, что разработанная система с некоторыми модификациями может быть положена в основу систем распознавания больших словарей, слитной речи и систем понимания речи.

В заключение автор выражает признательность В.М. Величко, В.Г. Лебедеву и А.Н. Манохину за ряд ценных замечаний, сделанных в процессе дискуссий при выполнении данной работы.

Л и т е р а т у р а

1. ЛОЗОВСКИЙ В.С. Аппроксимация отклика системы в π -пространстве формантным анализом речи. - В кн.: Вычислительные системы. Вып. 37. Новосибирск, 1969, с. 22-37.

2. ATAL B.S., HANAUER S.L. Speech analysis and Synthesis by linear prediction of Speech waveform. - "J. Acoust. Soc. Amer.", 1971, v. 50, p. 637-655.

3. ISHIKAWA I., NAKANO Y., NAKATA K. Evaluation of Various parameter sets in spoken digits recognition. - "IEEE Trans. Audio Electroacoust.", 1973, v. AU-21, N 3, June, p. 202-209.

4. ITAKURA F. Minimum prediction residual principle applied to speech recognition. - "IEEE Trans. Acoust., Speech, Signal Processing", 1975, v. ASSP-23, N 1, February, p. 67-72.

5. MARKEL J.D., GRAY A.H. On autocorrelation equation as applied to speech analysis. - "IEEE Trans. Audio Electroacoust.", 1973, v. AU-21, N 2, April, p. 69-79.

6. WHITE G.M., NEELY R.B. Speech recognition experiments with linear prediction, bandpass filtering and dynamic programming. - "IEEE Trans. Acoust., Speech, Signal Processing", 1976, v. ASSP-24, N 2, April, p. 183-188.

7. БОКС Дж., ДЖЕНИНС Г. Анализ временных рядов, прогноз и управление. Т. I, М., "Мир", 1974.

8. ГРЕНАНДЕР У., СЕГЕ Г. Теплицевы формы и их приложения. М., ИЛ, 1961.

9. ВЕЛИЧКО В.М., ЗАГОРТУЙКО Н.Г. Автоматическое распознавание ограниченного набора устных команд. - В кн.: Вычислительные системы. Вып. 36. Новосибирск, 1969, с. 101-110.

10. ВЕЛИЧКО В.М. Обработка речи на ЭВМ методом ассоциативного кодирования. - В кн.: Вычислительные системы. Вып. 62. (Ассоциативное кодирование.) Новосибирск, 1975, с. 38-48.

11. ФЛАНАГАН Д.Л. Анализ, синтез и восприятие речи. М., "Связь", 1968.

Поступила в ред.-изд. отд.
14 сентября 1978 года