

ИЗСЛЕДВАНЕ „ГЛАСНИ И СЪГЛАСНИ“⁽¹⁾

Мирослав Янакиев

Софийски университет „Св. Климент Охридски“

Целта е да отговорим най-напред на въпроса дали в български текст гласните букви (а, е, и, о, у, ъ, ю, я) се срещат по-често от съгласните букви (б, в, г, д, ж, з, й, к, л, м, н, п, р, с, т, ф, х, ц, ч, ш, щ; ъ означава „мекост“ при изговора на предходната съгласна и затова изобщо няма да го броим).

Можем да се опитаме да получим отговор на въпроса, като се заемем да броим съгласните и гласните букви в разни български текстове.

Кумулативно изследване

Този начин на търсене на отговор на въпроса (можем да го наречем индуктивен, но за предпочитане е да го наричаме кумулативен или натрупващ) изисква да се съобразяваме с правилата за математико-статистическа обработка на резултатите от броенето. Според тези правила трябва от големи набори числа („вариационни редове“) да извлечем две числа, които да послужат като отговор на въпроса. Едното от тези числа е „средното аритметично“, а другото – „средното квадратично отклонение“, оценява колко е вероятността да се окаже, че сме сбъркали, като сме приели за верен отговора, изразен чрез първото число.

Преди да се появят електронните калкулатори (елките, както у нас ги наричат), математико-статистическата обработка на много числа, свеждането им до две – „средно аритметично“ и „средно квадратично отклонение“ – изискваше много, макар и прости, изчисления, в които естествено се прокрадваха изчислителни грешки. Сега филолог, който разполага с подходящ калкулатор, може за две-три минути да получи „средното аритметично“ и „средното квадратично отклонение“ от „вариационни редове“, съдържащи стотина числа.

[...]

Най-простият начин да се даде отговор на въпроса по-често ли се срещат в български текст съгласни букви чрез кумулативно изследване, след като имаме калкулатор с вградена статистика, изисква най-напред да подберем текст. Нека текстът да бъде началото на стихотворението на Хр. Ботев „Хаджи Димитър“.

След това трябва да почнем да отброяваме *проби* от по десет букви и да номерираме пробите. Нека наричаме всяка такава проба *декаграма* (дека – „десет“, и грама – „буква“; уговорихме се „б“ да не броим!). И така, първата декаграма за изследването ни е: *жив е той жив*. Втората: *е там на балк*. Третата: *ана потънал*. И т. н.

Като насъберем десет декаграми (може да не бъдат непременно наред една след друга), преброяваме съгласните в първата декаграма и броя им „внасяме“ чрез натискане на клавиша, край който има надпис „х“, като предварително превключим калкулатора по указанията за използването му в режим „статистика“, което личи на дисплея му по появата на ST, STAT, σ , Σ (или на някакъв друг знак). След това внасяме броя на съгласните от втората декаграма, от третата и т.н. до десетата.

[...]

За да получим средното аритметично („хикс черта“ го наричат статистици-те), се натиска клавиш, край който има надпис „х“ с чертичка отгоре [x].

С помощта на калкулатора, чрез изчисляване на „средното квадратично отклонение“ (отбелязвано с гръцката буква „ σ “ с индекс „n – 1“ [σ_{n-1}] и затова наричано от статистиците „сигма ен минус едно“, [...] се преценява дали кумулацията (натрупването) на данните стига, за да се отговори на въпроса.

В нашия случай това става така: след като са внесени количествата на съгласните в първите десет декаграми на „Хаджи Димитър“ (6; 6; 5; 6; 5; 6; 5; 7; 6; 5), „хикс черта“ се оказва равно на 5 и 7 десети – на дисплея „5.7“ (в калкулаторите и на екраните на компютрите вместо „десетична запетая“ се използва „десетична точка“). От това следва, че средно взето, в десетте декаграми (проби) съгласните букви са повече от гласните.

Остава обаче да се провери чрез „средното квадратично отклонение“ доколко е вероятно да сбъркаме, ако приемем, че *изобищо* в текст на български език не можем да очакваме да се появят по-малко съгласни букви, отколкото гласни.

Въз основа на математико-статистическа закономерност, популярна под названието „закон за трите сигми“, може да се твърди, че ако 1) разделим „средното квадратично отклонение“ на квадратен корен от броя на пробите (в случая на „корен квадратен“ от „10“); 2) умножим полученото число на три; 3) полученото произведение извадим от средното аритметично и резултатът (ще го наречем „хикс черта минимум“) е по-голям от 5, то вероятността да попаднем случай на десетица декаграми, в които „хикс черта“ да е по-малко от 5 (т.е. в която съгласните букви да са по-малко от гласните), е практически равна на нула (по-точно, по-малка от 0,0003).

В нашия случай „средното квадратично отклонение“ е равно на 0,67. Значи, „хикс черта минимум“ е равно на:

$$5,7 - 3 \cdot 0,67/\sqrt{10} = 5,7 - 0,64 = 5,06$$

[...]

Следователно можем да смятаме, че отговорът на въпроса е намерен: практически невъзможно е да попаднем на български текст от поне 100 букви със съгласни повече от гласните.

[...]

Защо е уместно названието *кумулятивно* за изследване от типа на приведеното? Защото, ако започнем да изчисляваме средната аритметична, средното квадратично отклонение и минималната средна аритметична отначало за две проби, после за 3, после за 4 и т. н., ще видим, че за по-малки извадки (= вариационни редове) минималната средна аритметична се оказва недостатъчно стабилна над 5, т. е. трябва да натрупаме повече проби в извадката, за да докажем верността на отговора си. Или по-точно: като увеличаваме (кумулираме) броя на текстовите проби, ние установяваме каква е минималната дължина на текста, при която да е вярно твърдението: „Съгласните букви в такъв по големина български текст практически е невъзможно да бъдат по-малко от гласните.“

Направеното уточняване е много важно. Филологът изследвач трябва внимателно да формулира твърденията си, за да не се стигне до справедливи възражения.

Можем да си представим досетлив човек, който да разсъждава така: „Аз съм в състояние да посоча текст, който да се състои почти само от гласни букви – в крайните западни области на България е популярна остроумната диалектна фраза „Я ю яя, яя и ю до яя“, т.е. „Аз я яздох, яздох и я дояздох“, като отговор на въпрос какво се е случило с паднала кобила. А текст, който да се състои само или даже почти само от съгласни, не мога да си представя. Следователно не мога да приема, че в българския език няма текстове, състоящи се от повече гласни и по-малко съгласни букви!“.

Цитираната фраза „Я ю яя...“, записана, се състои от 11 гласни и само една съгласна буква (д), но опитайте се да съчинявате смислени български текстове от по 100 букви, в които да има поне 51 гласни, а останалите букви да са съгласни. Даже и да ви се удаде, ясно е, че едва ли ще съумеете да населите българската езикова практика с толкова много такива текстове, че средно взето на всеки 10 000 български текста, съставени от по 100 букви, да се появят поне по три, които, записани, да съдържат поне по 51 гласни. По „закона за трите сигми“ резултатът от кумулативното изследване на десетте декаграми от „Хаджи Димитър“ твърди, че именно това е невъзможно.

И така, въпросът „по-чести ли са съгласните букви от гласните в български текст“, има вече отговор.

Наистина, можем да чуем съмнения дали има някаква полза от подобни изследвания. Но филологът няма защо да страда от угризения на съвестта. Видяхме, че с помощта на джобното калкулаторче и не особено трудоемко броене той може да... задоволи любопитството (свое и на много други хора, не само българи) с досега неизвестна информация.

Въпросът „Каква полза?“ не заслужава по същество внимание, тъй като огромен брой младежи (и не само младежи) в света играят така наречените „електронни игри“, вградени в калкулаторчета, и се оправдават с това, че игрите ги „тренират“, дават им възможност да развиват или бързината на реакциите си, или умствените си способности.

На изследването „Гласни и съгласни“ може да се гледа като на тренировъчна игра, която подготвя за решаване на по-сложни филологически задачи с помощта на електрониката.

Ето една такава задача.

Изследване: има ли алитерация на „е“ в „Хаджи Димитър“ на Христо Ботев.

В своите лекции по история на българската литература проф. Боян Пенев е твърдял, че в четири строфи от стихотворението на Хр. Ботев „Хаджи Димитър“ има алитерация на „е“, т.е. гласната „е“ се е срещала неслучайно често. Много наши учители в уроците си по литература, като обясняват какво е алитерация или звукопис, сочат тези строфи. Ето ги (гласната „е“ нарочно съм подчертал и с номер и наклонени черти съм разделил текста на декаграми):

Настане веч¹/ер, месец изг²/рее,
звезди о³/бсипят свод⁴/а небесен.
Го⁵/ра зашуми, ве⁶/тер повее –
Ба⁷/лканът пее х⁸/айдушка пес⁹/ен.

И самодив¹⁰/и в бела прем¹¹/ена,
чудни, пр¹²/екрасни, пес¹³/ен поемнат,
т¹⁴/ихо нагазят¹⁵/ трева зелен¹⁶/а
и при юнака¹⁷/ дойдат, та се¹⁸/днат.

Една му¹⁹/ с билки рана²⁰/та върже,
дру²¹/га го пръсне²²/ с вода студ²³/на,
тре²⁴/та го в²⁴/ уста целуне²⁵/ бърже,
а той я²⁶/ гледа мила, з²⁷/асмена.

И плес²⁸/снат с ръце, п²⁹/а се прегърн³⁰/ат,
и с песни х³¹/въркнат те в н³²/есесата.
Лет³³/ят и п³⁴/ят, дор³⁴/де осъмнат
и³⁵// търсят духа/ на Караджат/а.

Да изчислим средната аритметична и средното квадратично отклонение за гласната буква „е“ в 35-те декаграми от „Хаджи Димитър“, цитирани по-горе (четирите строфи се състоят от пълни 37 декаграми, но последните две ще пренебрегнем – в тях няма „е“; значи, правим отстъпка на хипотезата на проф. Б. Пенев).

Калкулаторът ни дава ср. аритм. = 1.42857... $\approx$ 1.43, ср. квадр. отклонение = 0.8840 $\approx$ 0.88 и мин. ср. аритметична = 0.980... $\approx$ 0.98. Общо калкулаторът наброява в 35-те декаграми 50 гласни букви „е“ [По-нататък М. Янакиев обозначава сумата на числата във вариационния ред със Σx .] [...].

За да установим вярна ли е хипотезата на Б. Пенев, трябва да проверим каква е вероятността да се появят случайно в български текст с размер 35 декаграми поне толкова гласни букви „е“. Ако тази вероятност се окаже много малка, ще имаме основание да приемем, че струпването на 50 гласни букви „е“ в 35-те декаграми от Ботевото стихотворение не е случайно, т.е. че Ботев специално се е старал да подбира думи, съдържащи „е“-та, и Б. Пенев е прав.

За да определим тази вероятност, математическата статистика иска от нас да разполагаме с оценки за средното аритметично и за средното квадратично отклонение, получени от „среден“ (какъв да е!) български текст. Засега с такива оценки не разполагаме. Това е една от задачите, които българските филолози ще решават с помощта на компютрите. Досега тя не е получила задоволително решение, понеже, без да използваме компютри, за да я решим, трябва да работим десетилетия.

Но (временно!) можем да се опитаме да я решим, като повикаме на помощ логиката: ще приемем, че нужните ни оценки за средното аритметично и средното квадратично отклонение имаме право да получим, като изследваме български текст, за който да сме сигурни, че не е съчиняван с цел да се постигне поетично (звукописно) въздействие.

Като такъв текст ще използваме случаен откъс от съобщение, поместено във в. „Орбита“ (бр. 2/1987, с. 13). Ето този текст (пак разделен на декаграми):

Себестойно¹/стга на слън²/чевите елем³/енти бързо н⁴/амалаява: за п⁵/оследните д⁶/есет години⁷/ разходите п⁸/о създаване⁹/то на слънче¹⁰/ви батерии с¹¹/а намалели н¹²/ад три пъти (з¹³/а единица мо¹⁴/щност). Това д¹⁵/ава основан¹⁶/ие на специа¹⁷/листите да п¹⁸/рогнозират¹⁹/, че в близко б²⁰/ъдеще слънч²¹/евите елем²²/енти ще стана²³/т важни конк²⁴/урентоспос²⁵/обни източн²⁶/ици на енерг²⁷/ия. През посл²⁸/едните дес²⁹/т-дванайсет³⁰/ години силн³¹/о нарасна ин³²/тересът към³³/ използван³⁴/ето на възст³⁵/ановяемите³⁶/ енергетичн³⁷/и ресурси на³⁸/ океана.

Калкулаторната обработка на вариационния ред от количествата на буквата „е“ в 35 декаграми от откъса дава $\Sigma x = 41$. С 9 по-малко са „е“-тата в откъса

от „е“-тата в 35-те декаграми от „Хаджи Димитър“. Като че Б. Пенев е прав...

Но да не прибързваме! Полученото средно аритметично ($\bar{x}$), приблизително равно на 1,17, и средното квадратично отклонение (σ), приблизително равно на 1,22, за „е“-тата в 35-те декаграми на в. „Орбита“, ако спазваме изискванията на математическата статистика, *опровергават* твърдението на Б. Пенев. И ето защо.

В средите на математическата статистика се приема така нареченият „закон за двете сигми“. Според този закон, ако разликата между средното аритметично на оценявания текст (в случая – откъса от „Хаджи Димитър“) и средното аритметично на случайно подбрения текст (в случая – текста от в. „Орбита“) е по-малко от $2\sigma/\sqrt{n}$ (в случая $\sigma = 1,22$, а $n = 35$, поради което $2\sigma/\sqrt{35} = 0,41$), тя не е толкова голяма, че да я смятаме за неслучайна (в случая тази разлика е $1,43 - 1,17 = 0,26$, т. е. по-малка от 0,41). Или, както се изразяват в стила на математическата статистика, „нулевата хипотеза не може да бъде отхвърлена“. „Нулева хипотеза“ наричат предположението, че изчислената разлика е чисто *случайна*!

И така, математическата статистика не потвърждава хипотезата на Б. Пенев. Значи ли това, че Б. Пенев не е прав?

Разумният филолог и сега не бива да прибързва със заключение. Трябва да се постараме да потърсим „рационалното зърно“ в твърдението на нашия учен. Човешката глава е устройство, което в работата си също се подчинява на законите на математическата статистика.

Можем да допуснем, че Б. Пенев, като е говорел за „е“-алитерация, е имал предвид акцентуваните „е“-та. Вариационният ред, състоящ се от 35-те декаграми от „Хаджи Димитър“, само за акцентуваните „е“-та дава $\bar{x} = 0,688$ и $\sigma = 0,530$. Вариационният ред, съставен от 35-те декаграми от в. „Орбита“, за акцентуваните „е“-та дава $\bar{x} = 0,286$ и за $\sigma = 0,519$. Разликата между средните аритметични на „Хаджи Димитър“ и на „Орбита“ ($= 0,402$) е повече от 4,55 пъти по-голяма от $\sigma_{(„Орбита“)} / \sqrt{35} = 0,088$. А достатъчно е тази разлика да е три пъти по-голяма от 0,088, за да приемат специалистите по математическа статистика „нулевата хипотеза“ за опровергана („законът за трите сигми“).

Значи за акцентуваните „е“-та (само за тях) в откъса от „Хаджи Димитър“ Б. Пенев е прав – те са неслучайно повече, отколкото в специално необогатяван от автора с акцентувани „е“-та български текст.

И все пак скептиците още могат да се съмняват: може би по-честата употреба на акцентувано „е“ е характерна изобщо за българската поезия. Малко работа с калкулатора и съмнението е отстранено: в първите 35 декаграми от Ботевото стихотворение „Обесването на Васил Левски“ средната честота на акцентуваното „е“ в декаграма е 0,20, т.е. почти като във в. „Орбита“ (0,29).

А в останалата част от „Хаджи Димитър“?

В разглеждания откъс от „Хаджи Димитър“ между третата и четвъртата строфа не беше включена една строфа, в която Хаджи Димитър се обръща към самодивата:

– Кажи ми, сест¹/ро, де ѝ Карад²/жата?

Де е и мо³/йта верна др⁴/ужина?

Кажи м⁵/и, па ми вземи⁶/душата!

Аз ис⁷/кам, сестро, т⁸/ук да загина.⁹

Основанието е, че в „пряката реч“ поетите се стремят обикновено да спазват нормите на естествения разговор, а алитерацията изисква подбор на изразните средства, които е вероятно да отклони текста от разговорните норми. Проверката на броя на акцентуваните „е“-та в цитираната строфа, като че потвърждава валидността на това основание – в 9-те декаграми акцентуваното „е“ се появява 5 пъти, значи средната често е 0,56. В цитираните по-рано 4 строфи има 37 декаграми, в които акцентуваното „е“, се появява 24 пъти, значи средната честота е 0,65.

Но 0,56 е много по-близо до 0,65, отколкото до средната честота на появяване на акцентувано „е“ във в. „Орбита“ и в „Обесването на Левски“ ($0,20 \div 0,30$). И даже да включим строфата с „пряката реч“ в алитериранията част на от „Хаджи Димитър“, получената средна честота на акцентуваното „е“ в общо 46-те декаграми $x = 0,63$ остава достатъчно по-голяма от $x \approx 0,20 \div 0,30$, за да остане в сила и твърдението, че всичките 5 строфи от „Хаджи Димитър“ се характеризират с наличие на неслучайно честа поява на акцентувано „е“.

А останалата част от „Хаджи Димитър“?

Ако обединим началото (до „Настане вечер...“) и края – от „Но съмна вече...“ нататък, ще получим вариационен ред от появите на акцентувано „е“ в 65 декаграми с $x = 0,46$ ($\sigma = 0,68$). Да обединим, от друга страна, в един вариационен ред появите на акцентувано „е“ в 38-те декаграми от в. „Орбита“ и появите му в 48-те декаграми от „Обесването на Васил Левски“ – получаваме $x = 0,28$ и $\sigma = 0,52$, откъдето средната погрешност е равна на $\sigma/\sqrt{86} = 0,056$. За да оценим колко далече е 0,46 от 0,28, трябва да измерим разликата, като използваме за единица мярка средната погрешност: $(0,46 - 0,28)/0,056 = 3,21$. Тъй като 3,21 е число, по-голямо от 3, между честотите на поява на акцентуваното „е“ в неалитериранията част на „Хаджи Димитър“ и средно в българската езикова практика също има значима разлика – излиза, че в цялото стихотворение „Хаджи Димитър“ акцентуваното „е“ се появява по-често, отколкото може да се очаква, ако както в „Обесването на Васил Левски“, авторът не е извършвал специален подбор на изрази, съдържащи акцентувано „е“.

Възниква въпрос: може ли да се смята, че вътре в самото стихотворение „Хаджи Димитър“ има два типа обогатяване на текста с акцентувано „е“.

Б. Пенев говори за звукопис само в посочените четири строфи. А може би честите акцентувани „е“-та са разсеяни случайно из цялото стихотворение и би трябвало да говорим за изцяло „е“-алитериран текст. Б. Пенев просто не е усетил повсеместността на алитерацията и трябва следователно още веднъж да бъде уточнен.

Математическата статистика е в състояние да даде отговор и на този въпрос, но без електронната изчислителна техника намирането на отговори от такъв тип изисква тежък изчислителен труд, работа с многоцифрени числа, логаритмуване, умножаване на многозначни логаритми, антилогаритмуване, работа, която носи рискове да бъркаме, без да можем лесно да откриваме грешките си.

[...]

Критерий на различието „хи-квадрат“

Така се нарича математико-статистическа процедура, предложена от К. Пирсън (С. Pearson). В резултат от тази процедура се получава число, наричано „хи-квадрат“ [χ^2]. „Хи“ е името на гръцката буква, която е първа в гръцката по произход дума „характеристика“, а „квадрат“ напомня, че числото е сума от *квадрати* на разлики.

По големината на това число се съди каква е вероятността различието между два вариационни реда да е само случайно.

Когато се съпоставят два вариационни реда, извлечени от конкретни наблюдения (наричат такива редове *емпирични* – старогръцката морфема „емпеир“, произнасяна по новогръцки „емпир“, съответства на българското „получен в резултат от наблюдаване на опити“), числото „хи-квадрат“ се получава по твърде сложна формула, която е обоснована в учебните помагала по математическа статистика (вж. например УРБАХ 1964, с. 228 – 234). Ще обясня работата по тази формула чрез пример.

Двата (емпирични) вариационни реда ще бъдат: първият, съставен от количествата на появите на акцентувано „е“ в 65-те (по предположение слабо алитерирани) декаграми в началото и в края на „Хаджи Димитър“; вторият, съставен от количествата на появите на акцентувано „е“ в 46-те силно алитерирани декаграми (пет строфи в средата на „Хаджи Димитър“ – от „Настане вечер“ до „духа на Караджата“). Вж. таблица 1.

Таблица 1. Слабо и силно е-алитерирани декаграми в „Хаджи Димитър“

Алитерация	Общо декаграми	Количество декаграми с по Е на брой акцентувани „е“-та				
		0	1	2	3	(Е)
слаба	(А) 65	41	19	4	1	(С)
силна	(В) 46	18	27	1	–	(D)

А сега формулата:

$$\chi^2\text{-квадрат} = \frac{(CB - DA)^2}{AB(C + D)}$$

Да заменим в тази формула символите с числа:

$$\chi^2 = \frac{(41 \times 46 - 18 \times 65)^2}{65 \times 46 \times (41 + 18)} + \frac{(19 \times 46 - 27 \times 65)^2}{65 \times 46 \times (19 + 27)} + \frac{(5 \times 46 - 1 \times 65)^2}{65 \times 46 \times (5 + 1)} =$$

$$= 2,906048410 + 5,643165625 + 1,517558528 = 10,06677256.$$

[...]

Според математическата статистика, ако полученото по използваната формула χ^2 е по-голямо от 6,65 (числата, подобни на това число, са дадени у УРБАХ 1964, с. 395, Табл. XIX), вероятността „различието между двата сравнявани (емпирични) вариационни реда да е случайно“, е по-малка от 0,01, т.е. тя се смята за толкова малка, че хипотезата „различието е случайно“ трябва да се отхвърли. Тъй като в нашия случай $\chi^2 = 10,1$ е значително по-голямо от 6,63, математическата статистика ни дава право да твърдим, че „Хаджи Димитър“ по отношение на алитернацията на акцентуваното „е“ се състои от две в различна степен алитерирани части: едната („рамкова“), обхваща началото и края – слабо алитерирана; другата („централна“) – силно алитерирана.

Ето как снабденият с инструментария на плотометрията филолог може да докаже, че в един художествен текст обективно съществуват незабележими за „простото око“ особености.

NOTES/БЕЛЕЖКИ

1. През 1988 г. Мирослав Янакиев издава малката книжка (152 с.) „Електрониката в помощ на учителя филолог“. Народна просвета. София, 1988. Авторът предостави в издателството текстови файлове и тайно се гордееше, че това е първата филологическа книга, набрана в електронен вид. Но това му изигра лоша шега. При липсата на опит (и в издателството, и у автора) при електронно издаване на текст, наборът е изпълнен с много грешки – размествания, замяна на символи, някои символи (например √) дори са изчезнали. При това, когато издателството предостави коректура (само за 24 часа!), М. Янакиев не беше в страната и не

можа да я види. Изобщо, изданието добре илюстрира пълния разпад в българското книгоиздаване от 80-те години. Днес тази книга е достъпна на страниците, посветени на проф. Мирослав Янакиев – miryan.org.

М. Янакиев е насочил книгата си към будни ученици и студенти, но особено към онези стотици учители по български език и литература, които в Софийския университет са натрупали повече или по-малко опит в областта на глотометричните изследвания. Книгата е изградена в модерния сега жанр „How To“ – такъв подход тогава беше за нас новост. Информацията от областта на филологията и математическата статистика авторът разяснява с подробни инструкции как да се събират данни и как да се обработват с помощта на достъпните тогава електронни устройства – от „обикновени“ калкулатори до 8-битовите компютри „Правец“, с които българските училища бяха пълни, но които обикновено стояха заключени. Ясно е, че днес повечето от тези устройства могат да бъдат видени само в някой музей на електрониката. И ако авторът на тази книжка не беше Мирослав Янакиев, тя отдавна щеше да е остаряла, да представлява само исторически интерес. Но авторът все пак е Мирослав Янакиев – човек забележителен и като преподавател, и като теоретик. Малка част от идеите, изложени в тази книга, предимно свързани с развитието на техниката, са осъществени вече или поне са правени опити да се осъществят.

По-голямата част от идеите обаче все още чакат да ги проумеем. А примерите? Примерите на М. Янакиев представляват всъщност малки, блестящи изследователски етюди. Например *Изследване: Има ли алитерация на „е“ в „Хаджи Димитър“ на Христо Ботев*. Както е известно, тази идея е на проф. Боян Пенев и оттогава насетне се преповтаря и развива във филологическите изследвания. Но аз не съм видял някой от тези автори досега да е използвал или поне да е споменал изследването на Янакиев. Затова имам основание да говоря за „непознатия Янакиев“. Надявам се, публикуваният тук откъс от книгата да запълни донякъде тази празнина. Освен това в него ясно личат характерните за проф. Мирослав Янакиев черти – преподавателският му талант с типичния за него „дискусионен“, колоквиален стил на изложение, както и забележителната му аналитична мисъл.

Личи и още нещо, което читателят филолог малко по-трудно осъзнава: личи как М. Янакиев минимизира труда по научното изследване. А това става само с много задълбочена теоретична подготовка и с много, много, много изследователски опит. С [...] са отбелязани съкратените части от текста – съкращавал съм предимно обясненията как се работи с несъществуващата вече електронна техника. Пак с квадратни скоби са въведени и някои означения, които М. Янакиев е разяснявал в предишен или в съкратен текст. Оправени са и техническите грешки, най-вече в математическите означения и изрази. Правописа и пунктуацията на автора съм се старал да запазя.

[Бел. ред. Този съпътстващ коментар, озаглавен „Непознатият Мирослав Янакиев. Двадесет години от кончината му“, и подготовката на материала на проф. М. Янакиев за публикуване в сп. „Чуждоезиково обучение“ са дело на *Александър Иванов*].

REFERENCES/ЛИТЕРАТУРА

Urbakh, V. Yu. (1964). *Biometricheskiye metody*. Moskva: Nauka [Урбах, В. Ю. (1964). *Биометрические методы*. Москва: Наука]

STUDY “VOWELS AND CONSONANTS”

Prof. Miroslav Yanakiev
