

DIE DARSTELLBARKEIT VON WORTFORM-MEHRDEUTIGKEITEN
IM BULGARISCHEN

Jurgen Kunze, Alexandar Ljudskanov

Im folgenden soll gezeigt werden, wie Wortform-Mehrdeutigkeiten in einer bestimmten Weise dargestellt werden können. Zunächst sind aber einige elementare Begriffe zu erklären.

Unter einer Wortform a verstehen wir eine Einheit, die in einem schriftlichen Text von zwei Leerstellen begrenzt wird, ohne selbst eine solche zu enthalten. Es spielt dabei keine Rolle, ob es eine Wortform a' gibt, die als Zeichenfolge in a echt enthalten ist: *прави* und *правила* sind zwei verschiedene Wortformen, die zunächst keinerlei Zusammenhang aufweisen. Wir beschränken unseren allgemeinen Rahmen auf solche Sprachen, für die der so skizzierte Begriff „Wortform“ adäquat ist. Die Menge A aller Wortformen einer Sprache L wird das Vokabular von L genannt, und jeder Satz L ist eine Folge über A . Den Begriff „(korrekter) Satz“ kann man natürlich verschieden festlegen, je nach dem welchen Maßstab der Korrektheit man verwendet. Wir wollen uns hier auf grammatisch korrekte Sätze beschränken. Das später verwendete Modell hängt von dieser Voraussetzung nicht ab.

Viele Wortformen weisen irgendwelche Mehrdeutigkeiten auf. Derartige Wortformen nennt man auch Homographie. Ein Beispiel dafür ist die Wortform *прави*, die sowohl Verbform als auch Adjektiv sein kann, was der folgende Satz zeigt:

Той прави прави линии.

Kein Homograph ist z. B. die Wortform *аз*. Die Mehrdeutigkeit von *прави* ist mit den herkömmlichen grammatischen Mitteln natürlich leicht beschreibbar. Aber eine derartige Beschreibung setzt eine gewisse Kenntnis über das System der Sprache L voraus. Wir wollen uns im folgenden jedoch auf den Standpunkt stellen, daß die Sprache L als Menge von Sätzen $L = \{\varphi_1, \varphi_2, \dots\}$ gegeben ist. Über die Struktur und den Inhalt dieser Sätze sowie über irgendwelche Verwandtschaften zwischen ihnen treffen wir keinerlei Voraussetzungen. Unser Ziel ist es, aus der Menge L gewisse Rückschlüsse auf das System zu ziehen, das dieser Menge zugrunde liegt (vgl. auch [5]). Wir werden zeigen, daß bereits aus der Kenntnis der Menge aller grammatisch korrekten Sätze des Bulgarischen notwendigerweise folgt, daß z. B. die Wortform *прави* mehrdeutig sein muß.

Wir skizzieren nun kurz den Aufbau des verwendeten Modells (vgl. auch [1, 2]). Es sei x eine Wortform. Mit $\mathfrak{L}(L, x)$ bezeichnen wir den Kon-

textbereich von x in L , d. h. die Menge aller Paare $[a, \beta]$ mit $a\beta \in L$. Statt $[a, \beta]$ schreiben wir $a^*\beta$. So ist z. B. $\kappa y n u x * c e k y p a$ ein Kontext von $h o b a$. Ist X eine nichtleere Menge von Wortformen (d. h. eine Menge mit $\emptyset \subset X \subset A$, $X = \{x^1, \dots, x^n\}$), so bezeichnen wir mit $\mathfrak{U}_1(L, X)$ die Menge aller Kontexte, die zu allen $\mathfrak{U}(L, x^r)$ für $1 \leq r \leq n$ gehören:

$$\mathfrak{U}_1(L, X) = \bigcap_{x \in X} \mathfrak{U}(L, x) = \mathfrak{U}(L, x^1) \cap \dots \cap \mathfrak{U}(L, x^n).$$

$H(L, x)$ (die Hülle von X in L) sei die Menge aller Wortformen $y \in A$ für die

$$\mathfrak{U}(L, y) \supset \mathfrak{U}_1(L, X)$$

gilt. Die Operation H besitzt folgende grundlegenden Eigenschaften: (vgl. [1] S. 433).

- (1) $X \subset H(L, X)$ für beliebige X ;
- (2) $H(L, H(L, X)) = H(L, X)$ für beliebige X ;
- (3) Wenn $X_1 \subset X_2$, so $H(L, X_1) \subset H(L, X_2)$.

Eine Menge X heißt in L abgeschlossen, wenn $H(L, X) = X$ ist. Wegen (1) ist X also genau dann abgeschlossen, wenn für jedes $y \in A$ aus $\mathfrak{U}(L, y) \supset \mathfrak{U}_1(L, X)$ die Beziehung $y \in X$ folgt. Im folgenden lassen wir das Argument L fort!

$X_1, \dots, X_n$ seien beliebige nichtleere Teilmengen von A .

Wir bezeichnen mit $\prod(X_1, \dots, X_n)$ das Cartesische Produkt von $X_1, \dots, X_n$:

$$\prod(X_1, \dots, X_n) = X_1 \times \dots \times X_n = \{x_1 \dots x_n, x_1 \in X_1, \dots, x_n \in X_n\}.$$

Die Folge $X_1, \dots, X_n$ heißt in L zulässig, wenn $\prod(X_1, \dots, X_n) \subset L$ ist (vgl. [1], S. 443).

Wir wenden uns nun den Begriffen Wortformenklasse und Wortformenklassensystem zu. Wortformenklassen sollen natürlich gewisse Teilmengen von A sein. Wir erläutern zunächst den Unterschied zu den Wortklassen im traditionellen Sinne. Die Wortklassen stellen eine ziemlich grobe Einteilung der Wortformen dar. In die Wortklasse Verb fallen z. B. alle Formen von Verben, obwohl die Formen ganz unterschiedliche Funktionen (und damit auch unterschiedliche Kontextbereiche) haben. Dies gilt etwa für die Wortformen $вижда$ und $виждам$. Die Verschiedenheit der Kontextbereiche hat ihre Ursache in den bei diesen Verbformen möglichen Subjekten:

$$\begin{array}{lll} \text{мой я} & \in \mathfrak{U}(\text{вижда}) & \in \mathfrak{U}(\text{виждам}) \\ \text{аз я} & \in \mathfrak{U}(\text{вижда}) & \in \mathfrak{U}(\text{виждам}) \end{array}$$

Beide Verbformen verhalten sich hinsichtlich der möglichen anderen Erweiterungen (Objekte usw.) jedoch gleich. Gerade umgekehrt ist das Verhältnis zwischen den Wortformen $виждам$ und $ходя$. Hier ist zwar kein Unterschied in den Subjekten vorhanden, wohl aber in den Objekten, da das eine Verb transitiv ist, das andere dagegen nicht:

$$\text{аз я} \quad \in \mathfrak{U}(\text{виждам}) \notin \mathfrak{U}(\text{ходя})$$

Man kann viele Unterschiede dieser und anderer Art finden. So haben z. B. Formen eines Verbs, die sich bei gleichem Numerus und gleicher Person

nur im Tempus unterscheiden, ebenfalls verschiedene Kontextbereiche, da die temporalen Adverbialbestimmungen gewissen Einschränkungen unterworfen sind. Unser Ziel ist es, die Einteilung des Vokabulars so zu wählen, daß alle Wortformen einer Wortformenklasse X eine ganz bestimmte grammatische Funktion gemeinsam haben. Die Wortformen *вижда* und *виждам* z. B. müssen also verschiedenen Wortformenklassen angehören, und zwar in dem Sinne, daß es keine Wortformenklasse X mit $\{\text{вижда}, \text{виждам}\} \subset X$ gibt.

Wir hatten weiter oben schon gesagt, daß wir die grammatischen Kategorien der betrachteten Sprache nicht verwenden wollen, sondern nur die Satzmenge L . Bei dieser Grundvoraussetzung entsteht natürlich sofort die Frage, wie man Wortformenklassen überhaupt bestimmen kann. Wir wollen diese Frage hier nur ganz kurz beantworten: wenn man sich erst einmal heuristisch Klarheit über den Begriff „Wortformenklasse“ verschafft und in verschiedenen Sprachen viele solcher Klassen X aufstellt, so macht man eine ganz überraschende Feststellung: Diese Mengen X sind stets abgeschlossen im oben erklärten Sinne. Dies scheint eine ganz fundamentale Eigenschaft der grammatischen Wortformenklassen zu sein. Es ist daher angebracht, diese Eigenschaft in Form einer axiomatischen Hypothese einzuführen. Diese Hypothese ist nicht nur eine „Verallgemeinerung aus der Realität“, sie läßt sich auch theoretisch begründen ([2], S. 445—454).

Für unsere späteren Erörterungen wird eine zweite Hypothese ebenfalls Bedeutung haben. Sie ist wesentlich leichter begründbar, erfordert aber zu ihrer Formulierung einige Vorbetrachtungen. Es sei $\varphi = a_1 \dots a_n$ ein Satz aus L . Jede der Wortformen a_v ($1 \leq v \leq n$) aktualisiert in φ eine gewisse Funktion. Zu jeder dieser Funktionen gehört nach unserer generellen Annahme eine Wortformenklasse, die wir mit X bezeichnen wollen. Wir betrachten die Folge $X_1, \dots, X_n$. Führt man dies für hinreichend viele Sätze einer konkreten Sprache durch, so ergibt sich immer wieder, daß diese Folge stets zulässig im oben definierten Sinne ist. Man kann daher (auch auf Grund entsprechender theoretischer Erwägungen) folgende Hypothese als Axiom aussprechen:

Bildet das Mengensystem $A = \{X, Y, \dots\}$ ein Wortformenklassensystem für eine Sprache L , so gibt es zu jedem Satz $\varphi = a_1 \dots a_n \in L$ eine Folge $X_1, \dots, X_n$ mit folgenden Eigenschaften:

(1) Es ist $a_v \in X$ für $1 \leq v \leq n$.

(2) Es ist X für $1 \leq v \leq n$.

(3) Es ist $\prod (X_1, \dots, X_n) \in L$, d. h. $X_1, \dots, X_n$ ist zulässig.

Dabei ist X diejenige Klasse, die einer der Funktionen entspricht, die a_v in φ aktualisiert. Selbstverständlich muß die Folge $X_1, \dots, X_n$ nicht eindeutig bestimmt sein; denn es gibt mehrdeutige Sätze φ , deren Mehrdeutigkeit durch unterschiedliche Wortklassenfolgen, die zu φ gehören, demonstriert werden kann. Ferner wird nicht behauptet, daß jede Folge, die (1) bis (3) erfüllt, die aktualisierten Funktionen der Wortformen repräsentiert. (Dieses Problem wird in [3] bei der Definition der Strukturtypen erläutert.) Der Inhalt der zweiten Hypothese wird durch die folgenden Beispiele noch klarer werden.

Betrachten wir nun einige bulgarische Sätze. Wir setzen dabei ein beliebiges System A voraus, das nur die beiden Hypothesen erfüllen soll. Es ist also keineswegs so, daß wir uns auf ein bestimmtes System beziehen,

etwa auf das, welches den üblichen Kategorien entspricht. Als ersten Satz wählen wir

q = той прави прави линии.

Auf Grund der zweiten Hypothese existiert zu φ eine zulässige Folge, die wir mit X_1, X_2, X_3, X_4 bezeichnen. Dabei gilt *той* $\in X_1$, *прави* $\in X_2$, *прави* $\in X_3$, *линии* $\in X_4$. Die vier Mengen X_i müssen außerdem abgeschlossen sein. Wir werden beweisen, daß notwendigerweise $X_2 = X_3$ sein muß, falls $\mathcal{A}$ reduziert ist. (Ein Wortformenklassensystem heißt reduziert, falls es kein System $\mathcal{I}' \subset \mathcal{I}$ gibt, so daß $\mathcal{I}'$ die beiden Hypothese erfüllt.) Daraus folgt zweierlei:

(a) *прави* ist eine mehrdeutige Wortform; denn es gibt zwei verschiedene Klassen, in denen diese Wortform vorkommt,

(b) *прави*₂ aktualisiert in q eine andere grammatische Funktion als *прави*₁.

Zum Beweis der Behauptung betrachten wir die Wortformen *чертае* und *дълги*. Man überzeugt sich leicht davon, daß

(1) $\mathcal{V}(\text{чертае}) \subset \mathcal{V}(\text{прави})$,

(2) $\mathcal{V}(\text{дълги}) \subset \mathcal{V}(\text{прави})$ und

(3) $\mathcal{V}(\text{прави}) \subset \mathcal{V}(\text{чертае}) \vee \mathcal{V}(\text{дълги})$ gilt.

Die erste Inklusion besagt: ist φ ein Satz, in dem *чертае* vorkommt, so kann man *чертае* durch *прави* ersetzen, ohne daß die grammatische Korrektheit verloren geht. Die zweite Inklusion besagt dasselbe für *дълги*. Die letzte Beziehung drückt schließlich folgendes aus: Ist φ ein Satz, in dem die Wortform *прави* vorkommt, so ist wenigstens einer der beiden Sätze grammatisch korrekt, die man erhält, wenn man *прави* durch *чертае*, bzw. durch *дълги* ersetzt. In allen Fällen darf die Substitution natürlich nur an einer Stelle vorgenommen werden: in unserem Satz q führt die (nicht erlaubte) gleichzeitige zweimalige Ersetzung von *прави* durch *чертае*, bzw. *дълги*, zu den abweichenden Ketten

той чертае чертае линии = φ^* ,

той дълги дълги линии = φ^{**}

Eine korrekte Anwendung von (3) auf q ergibt dagegen

той чертае прави линии = φ' .

*прави*₂ ist in q durch *чертае* substituierbar. φ' enthält *прави* ebenfalls. Dieses Mal ist *прави* durch *дълги* ersetzbar. Man erhält:

той чертае дълги линии = φ'' .

Aus (1) bis (2) ist nun folgende Aussage (4) beweisbar, aus der sich unsere Behauptung dann ergeben wird.

(4) Für jede Klasse $X \in \mathcal{I}$ gilt:

Wenn *чертае* $\in X$ oder *дълги* $\in X$, so *прави* $\in X$.

Für abgeschlossene Mengen X gilt: ist $x \in X$ und $\mathcal{V}(x) \subset \mathcal{V}(y)$, so ist $y \in X$. Denn wegen $x \in X$ ist offenbar $\mathcal{V}_1(X) \subset \mathcal{V}(x) \subset \mathcal{V}(y)$, also $\mathcal{V}(y) \subset \mathcal{V}_1(X)$. Daraus folgt nach der Definition der abgeschlossenen Mengen jedoch $y \in X$. Für $x = \text{чертае}$ und $y = \text{прави}$ erhalten wir daher aus (1): wenn *чертае* $\in X$, so *прави* $\in X$. Für $x = \text{дълги}$ und $y = \text{прави}$ ergibt sich aus (2): wenn *дълги* $\in X$, so *прави* $\in X$. Damit ist die Aussage (4) bewiesen. Daraus läßt sich nun unsere Behauptung $X_2 = X_3$ gewinnen. Wir müssen dabei allerdings einiges voraussetzen, dessen Herleitung hier zu weit führen würde (vgl. [3], [4]). Die Aussage (4) läßt sich bei Berücksichtigung der Tatsache, daß *чертае* und *дълги* initiale

Wortformen sind (vgl. [2]), dazu verwenden, daß man folgendes zeigt: ist das System I reduziert, so gilt für jede Folge $X_1, X_2, X_3, X_4, X_5 \in A$, folgendes:

(5) es ist $\varphi'' \in \prod (X_1, X_2, X_3, X_4)$ genau dann, wenn $\varphi \in \prod (X_1, X_2, X_3, X_4)$ ist.

Das ergibt: wegen $\varphi'' \in \prod (X_1, X_2, X_3, X_4)$ ist $\text{чептае} \in X_2$ (und daher $\text{права} \in X_2$). Da φ^* kein Satz ist, muß $\text{чептае} \in X_3$ sein. Daraus folgt sofort $X_2 \neq X_3$. Wegen $\text{долгу} \in X_3$ ist jedoch $\text{права} \in X_2$.

Wir demonstrieren nun noch die Mehrdeutigkeiten einiger anderer Wortformen. Wir werden uns dabei jedoch kürzer fassen, indem wir jeweils nur die entsprechenden Sätze und die sich dabei ergebenden Klassen aufführen. (Dabei bezeichnet X_j^i die Klasse, die der j -ten Stelle in φ^i entspricht.)

Die Wortform *пила* ist mehrdeutig:

In φ^1 *тя беше пила вино* liegt die Klasse $X_3^1 = \{\text{пила, донесла}, \}$ vor; in $\varphi^2 = \text{куних нова пила}$ dagegen die Klasse $X_3^2 = \{\text{пила, секира}, \}$.

Dies erkennt man so: Es gilt

$$\begin{aligned} \mathfrak{U}(\text{донесла}) &\subset \mathfrak{U}(\text{пила}), \\ \mathfrak{U}(\text{секира}) &\subset \mathfrak{U}(\text{пила}). \end{aligned}$$

Da nun für

$$\begin{aligned} \varphi_1' & \text{тя беше донесла вино,} \\ \varphi_2' &= \text{куних нова секира} \end{aligned}$$

aufgrund der zweiten Hypothese entsprechende Klassenfolgen existieren, folgt die Behauptung sofort, da die Klassen abgeschlossen sein sollen. Ferner ist $X_3^1 \neq X_3^2$, da wegen der Unkorrektheit von

куних нова секира

донесла $\in X_3^2$ ist.

Bei beiden Beispielen würde eine ausführliche Darlegung nur die beiden oben genannten Hypothesen verwenden. Es ist jedoch niemals ein Rückgang auf die üblichen grammatischen Kategorien nötig, die Argumentationen sind davon völlig unabhängig.

Durch derartige Betrachtungen kann man sich sukzessive gewisse Klassen für eine vorgegebene Sprache L verschaffen. Es ist jedoch nicht immer so, daß dieses Verfahren mit einem eindeutigen Ergebnis endet. In einigen Fällen hat man gewisse „Freiheiten“ in der Auffassung. Derartige Freiheiten entstehen vor allem dann, wenn das System der Kategorien durch die Satzmenge nicht eindeutig bestimmt ist. Wir erläutern dies an einem etwas eigenartig erscheinenden Beispiel aus dem Bulgarischen. Die Wortform *ни* ist mehrdeutig, was die folgenden Sätze zeigen:

ти ни виждаш
ти ни давах книгата.

Nach den vorangegangenen Betrachtungen ist es naheliegend, die Mehrdeutigkeit dieser Wortform auf ganz analoge Weise durch Betrachtung der Wortformen *зо* und *му* zu zeigen; denn die Einsetzung von *зо* im ersten Satz und von *му* im zweiten Satz führt wieder zu korrekten Sätzen, außerdem gilt jedoch auch

$$(*) \quad \begin{cases} \mathcal{Q}(zo) \subset \mathcal{Q}(ни), \\ \mathcal{Q}(му) \subset \mathcal{Q}(ни), \end{cases}$$

so daß man entsprechend auf

$$\begin{array}{l} \text{Klassen} \\ \text{und} \end{array} \quad \left\{ \begin{array}{l} zo, ни, ви, ги, \\ му, ни, ви, им, \dots \end{array} \right\}$$

käme. Aber das ist ein Fehlschluß! In Wirklichkeit gelten die Inklusionen (*) nämlich nicht, d. h. es gibt Kontexte, in denen *zo* nicht durch *ни* ersetzt werden kann:

ние zo виждаме ist korrekt,

ние ни виждаме dagegen nicht!

Damit ist der bisher beschrittene Weg der Demonstration von Mehrdeutigkeiten nicht mehr möglich, wir müssen etwas anders vorgehen. Wenn wir von der morphologischen Form des Verbs einmal absehen, so passen in den „Kontext“

Subjekt (Pronomen) Objekt (Pronomen) *вижд*—
folgende Paare (durch + angedeutet):

	<i>ме</i>	<i>те</i>	<i>зо</i>	<i>я</i>	<i>ни</i>	<i>ви</i>	<i>ги</i>	<i>се</i>	
<i>аз</i>		+	+	+	+	+	+	+	X_1
<i>ти</i>	+		+	+	+	+	+	+	X_2
<i>той</i>	+	+		+	+	+	+	+	X
<i>тя</i>	+	+	+		+	+	+	+	X
<i>ние</i>	+	+	+	+		+	+	+	X_3
<i>вие</i>	+	+	+	+	+		+	+	X_4
<i>те</i>	+	+	+	+	+	+		+	X

Diese Situation erlaubt nun zwei verschiedene Klasseneinteilungen für die Objektspronomen:

(A) Die Klassen sind die mit + versehenen Zeilen des Schemas:

$$\{те, зо, я, ни, ви, ги, се\} = X_1,$$

$$\{ме, зо, я, ни, ви, ги, се\} = X_2,$$

$$\{ме, те, зо, я, ни, ви, ги, се\} = X,$$

$$\{ме, те, зо, я, ви, ги, се\} = X_3,$$

$$\{ме, те, зо, я, ни, ги, се\} = X_4.$$

Nun ist unmittelbar zu erkennen, daß

$$X^* = X_1 \cup X_2 \cup X_3 \cup X_4$$

ist. Wenn man sich auf reduzierte Wortformenklassensysteme beschränkt ([2], S. 455), kann man diese Klasse fortlassen, so daß man auf vier Klassen kommt. (Diese Klassen seien Teil eines Wortformenklassensystems für das Bulgarische, das wir mit I_1 bezeichnen.) Dann „paßt“ jede dieser Klassen mit dem Subjektspronomen der 1. und 2. Pers. zusammen, dessen Akkusativ nicht in ihr enthalten ist. Die Subjektspronomen der 3. Pers. passen mit jeder Klasse zusammen. Wir stellen dies im folgenden Schema dar:

	X_1	X_2	X_3	X_4
<i>аз</i>	+			
<i>ти</i>		+		
<i>мой</i>	+	+	+	+
<i>тя</i>	+	+	+	+
<i>ние</i>			+	
<i>вие</i>				+
<i>те</i>	+	+	+	+

(B) Die Klassen haben folgende Gestalt:

$$\begin{aligned}\{ме, го, я, ги, се\} &= Y_1, \\ \{те, го, я, ги, се\} &= Y_2, \\ \{ни, го, я, ги, се\} &= Y_3, \\ \{ви, го, я, ги, се\} &= Y_4.\end{aligned}$$

Dieses Mal ergibt sich folgendes Schema für die Verträglichkeit mit dem Subjektspronomen:

	Y_1	Y_2	Y_3	Y_4
<i>аз</i>		+	+	+
<i>ти</i>	+		+	+
<i>мой</i>	+	+	+	+
<i>тя</i>	+	+	+	+
<i>ние</i>	+	+		+
<i>вие</i>	+	+	+	
<i>те</i>	+	+	+	+

Wir nehmen an, daß diese Klassen Teil eines Wortformenklassensystems $\mathcal{A}_2$ sind.

Entsprechende Betrachtungen kann man für den Dativ durchführen. Auch hierfür ist die Klasseneinteilung nicht eindeutig.

Schließlich kann man die Personalpronomina für Akkusativ und Dativ gemeinsam in ein System von Klassen einteilen. Es ergibt sich dann die Möglichkeit, eine Klasse $\{ни\}$ einzuführen, da es keine Wortform x mit $\mathcal{V}(ни) \subset \mathcal{V}(x)$ gibt, und diese Klasse ist die einzige, die *ни* enthält. Diese Klasse sei in einem Wortformenklassensystem $\mathcal{A}_3$ enthalten. Bei dieser Auffassung erscheint die Wortform nicht als mehrdeutig.

Fassen wir zusammen: die Wortform *ни* ist hinsichtlich ihrer Mehrdeutigkeit nicht eindeutig festgelegt. In $\mathcal{A}_1$ und $\mathcal{A}_2$ ist sie mehrdeutig, in $\mathcal{A}_3$ eindeutig. Der tiefere Grund dafür liegt darin, daß sie ein freies Homonym ist (vgl. [2], S. 423). Ein freies Homonym ist eine Wortform z , zu der keine Wortform y mit $\mathcal{V}(y) \subset \mathcal{V}(z)$ existiert, aber wenigstens eine Wortform x mit $\mathcal{V}(z) \subset \mathcal{V}(x) \neq \emptyset$ und $\mathcal{V}(z) \subset \mathcal{V}(x)$. Wir wollen an Beispielen kurz illustrieren, daß dies für die Wortform $z = \text{гилт}$ gilt. Betrachten wir zunächst die Möglichkeit $\mathcal{V}(y) \subset \mathcal{V}(z) -$
 $y = \text{ви}$ kommt nicht in Frage, da

$ни\acute{e}\text{ виждаме} \in \mathcal{V}(y), \in \mathcal{V}(z);$
 $y = ce$ kommt nicht in Frage, da
 $ни\acute{e}^* \text{ виждаме} \in \mathcal{V}(y), \in \mathcal{V}(z);$
 $y = cu$ kommt nicht in Frage, da
 $ти\text{ спокоен} \in \mathcal{V}(y), \bar{\in} \mathcal{V}(z);$
 $y = me$ kommt nicht in Frage, da
 $^* \text{ са спокойни} \in \mathcal{V}(y), \bar{\bar{\in}} \mathcal{V}(z).$

Auf diese Weise kann man alle Pronomina ausschließen. (Andere Wortformen kommen aus syntaktischen Gründen ohnehin nicht in Frage.)

Als x kann man *те* wählen: Wegen

$той\text{ вижда} \in \mathcal{V}(x) \cap \mathcal{V}(z)$

und

$той^* \text{ дава книгата} \in \mathcal{V}(x), \in \mathcal{V}(z)$

leistet diese Wortform das Verlangte.

Damit haben wir gezeigt, daß die oben erwähnten Freiheiten bei der Konstruktion von Wortformenklassensystemen tatsächlich existieren. Genauer gesagt, die beiden Hypothesen erlauben für die bulgarische Wortform *ни* keine Antwort auf die Frage der Mehrdeutigkeit. „In Wirklichkeit“ ist sie natürlich mehrdeutig. Wir wollten nur die Möglichkeiten und die Grenzen zeigen, die durch die beiden Hypothesen gegeben sind. Durch andere Hypothesen kann die Mehrdeutigkeit von *ни* sehr einfach gezeigt werden.

LITERATURVERZEICHNIS

1. Kunze, J. Versuch eines objektivierten Grammatikmodells I. — Z. für Phonetik, Sprachwiss. und Kommunikationsforschung, **20**, 1967, No. 5/6, 415–448.
2. Kunze, J. B. II. Ibid., **21**, 1968, No. 5, 421–466.
3. Kunze, J., W. Priess. B. III. Ibid., **23**, 1970, No. 4, 347–378.
4. Kunze, J., W. Priess. Aufgaben und Möglichkeiten eines objektivierten Grammatikmodells. — Conference on Computational Linguistics, Balatonszabadi, September 1969 (Vortrag).
5. Ljudskanov, A. Is the Generally Accepted Strategy of Machine-Translation Research Optimal? — Mechanical Translation and Computational Linguistics, **11**, 1968, No 1/2, 14–22.

Eingegangen am 22. XI. 1971

ЕДИН НАЧИН ЗА ПРЕДСТАВЯНЕ НА ОМОНИМИЯТА НА СЛОВОФОРМИТЕ В БЪЛГАРСКИЯ ЕЗИК

Юрген Кунце, Александър Людсканов

(Резюме)

Накратко се излага модел, който осигурява формално описание на омонимията (омографията) в даден език L , и се дава съответно интерпретиране с материал от български език,

Проблемата се свежда към следното.

Традиционните начини на описание на омонимията предполагат познания за дадения L . За разлика от това в работата се излиза от предпоставката, че L е зададен само като множество вериги (изречения) и целта е, като се изхожда от него, да се дедуцират заключения за системата, която лежи в основата на този език, в това число и за омонимията в него.

Въвеждат се и се дефинират следните основни понятия и условия: контекстна област; множество от всички контексти за $\mathcal{V}(L, x^r)$; условията, при които едно множество X от словоформи е затворено в L , както и условията за допустимостта в него на вериги от вида $X_1 \dots X_n$; след това се дефинират понятията „клас словоформи“ и „система от класове на словоформи“, въвеждат се две аксиоматични хипотези и се показва как при дадена система от класове словоформи $\mathcal{I}$ субституирането на дадени „влизания“ в допустими вериги при условие да се запазва граматическата коректност отнася омонимичните словоформи към различни класове, когато $\mathcal{I}$ е редуцирана, с което се и решава въпросът за тяхното описание.

В работата се дава съответна формална обосновка, показва се, че предложеният метод позволява последователно еднозначно обособяване на класове в предварително зададени L и оставя известна „свобода“ на съждение при конструирането на системи от класове тогава, когато това условие не е налице. Това се показва чрез интерпретиране на фрагменти от системата на българските местоимения.

ОДИН ВОЗМОЖНЫЙ СПОСОБ ПРЕДСТАВЛЕНИЯ ОНОНИМИИ СЛОВОФОРМ В БОЛГАРСКОМ ЯЗЫКЕ

Юрген Кунце, Александр Людсканов

(Резюме)

Приводится сокращенное изложение модели, обеспечивающей формальное описание омонимии (омографии) в данном языке L , и дается интерпретация модели на материале болгарского языка. Проблема сводится к следующему.

Традиционные способы описания омонимии предполагают знание данного L . В отличие от этого работа исходит из предположения, что L задан только как множество цепочек (предложений), и цель заключается в дедуцировании заключений относительно системы, лежащей в основе этого языка, в том числе и в отношении омонимии в нем.

Вводятся и дефинируются следующие основные понятия и условия: контекстная область; множество всех контекстов для $\mathcal{V}(L, x^r)$; условия, при которых данное множество словоформ X замкнуто в L , как и условия допустимости в нем цепочек вида $X_1, \dots, X_n$; после этого дефинируются понятия „класс словоформ“ и „система классов словоформ“

вводятся две аксиоматические гипотезы и показывается, как при данной системе классов словоформ L субституирование данных „вхождений“ в допустимые цепочки при условии сохранения грамматической корректности относит омонимические словоформы к различным классам, когда L редуцирована, чем и решается вопрос их описания.

В работе дается соответствующее формальное обоснование, показывается, что предложенный метод позволяет последовательное однозначное обособление классов в предварительно заданных L и оставляет известную „свободу“ суждения при создании систем классов при отсутствии этого условия, что иллюстрируется интерпретированием фрагментов системы болгарских местоимений.