
**Language resource management —
Morpho-syntactic annotation framework
(MAF)**

*Gestion des ressources langagières — Cadre d'annotation
morphosyntaxique (MAF)*

iTeh Standards
(<https://standards.iteh.ai>)
Document Preview

ISO 24611:2012

<https://standards.iteh.ai/catalog/standards/iso/442a57b7-65de-4d5f-9c8f-34c401137730/iso-24611-2012>



iTeh Standards
(<https://standards.iteh.ai>)
Document Preview

ISO 24611:2012

<https://standards.iteh.ai/catalog/standards/iso/442a57b7-65de-4d5f-9c8f-34c401137730/iso-24611-2012>



COPYRIGHT PROTECTED DOCUMENT

© ISO 2012

All rights reserved. Unless otherwise specified, no part of this publication may be reproduced or utilized in any form or by any means, electronic or mechanical, including photocopying and microfilm, without permission in writing from either ISO at the address below or ISO's member body in the country of the requester.

ISO copyright office
Case postale 56 • CH-1211 Geneva 20
Tel. + 41 22 749 01 11
Fax + 41 22 749 09 47
E-mail copyright@iso.org
Web www.iso.org

Published in Switzerland

Contents

Page

Foreword	v
Introduction	vi
1 Scope	1
2 Normative references	1
3 Terms and definitions	1
4 The MAF meta-model	4
4.1 Overview	4
4.2 MAF Meta-model	4
5 Segmenting with tokens	6
5.1 General	6
5.2 Formal description: <token>	7
5.3 Embedding notation	7
5.4 Alternate representation for TEI based documents	8
5.5 Stand-off notation	9
5.6 Informative attributes	9
5.7 Completing the inline token notation	10
5.7.1 Joining tokens in embedded mode	10
5.7.2 Overlapping tokens	11
6 Word-forms as linguistic units	11
6.1 Formal description: <wordForm>	12
6.2 Token attachment	12
6.2.1 One token; one word-form	12
6.2.2 Several contiguous tokens; one word-form	12
6.2.3 Several discontinuous tokens; one word-form	13
6.2.4 Zero token; one word-form	13
6.2.5 One token; several word-forms	14
6.3 Referring to lexical entries	14
6.4 Compound word-forms	15
6.5 Identification of word-forms within a TEI-compliant document	15
7 Morpho-syntactic content	18
7.1 General	18
7.2 Using feature structures	18
7.3 Compact morpho-syntactic tags	18
7.4 FSR libraries	19
7.5 Designing tagsets	20
7.6 Formal description: <tagset>	22
8 Handling ambiguities	22
8.1 Word-form content ambiguities	22
8.2 Lexical Ambiguities	23
8.3 Structural ambiguities	23
8.3.1 Structural ambiguities with word-forms	23
8.3.2 Structural ambiguities with tokens	24
8.4 Simplified structuring variants	24
8.4.1 Non-ambiguous linear representation	24
8.4.2 Mixed linear and lattice representation	25
8.5 Expanding the simplified variants	26
8.5.1 Separating tokens and word-forms	26
8.5.2 Wrapping into local lattices	26

8.5.3	Merging local lattices	27
8.5.4	Removing <wfAlt>	28
8.6	Formal description: <wfAlt> and <fsm>	29
Annex A (informative) Encoded example using the MAF serialization.....		30
Annex B (normative) MAF specification		33
B.1	Elements	33
B.1.1	<dcs/>	33
B.1.2	<fsm>	34
B.1.3	<maf>	34
B.1.4	<tagset>	35
B.1.5	<token>	35
B.1.6	<transition>	36
B.1.7	<wfAlt>	36
B.1.8	<wordForm>	37
B.2	Model classes.....	38
B.3	Attribute classes	38
B.3.1	att.token.information	38
B.3.2	att.token.join.....	39
B.3.3	att.token.span.....	39
B.3.4	att.wordForm.content.....	39
B.3.5	att.wordForm.tokens	40
B.4	Macros	40
B.4.1	data.certainty.....	40
B.4.2	data.code	40
B.4.3	data.count.....	40
B.4.4	data.duration.w3c	41
B.4.5	data.enumerated	41
B.4.6	data.key.....	41
B.4.7	data.language.....	42
B.4.8	data.name	43
B.4.9	data.numeric.....	43
B.4.10	data.pointer	43
B.4.11	data.probability	44
B.4.12	data.temporal.w3c.....	44
B.4.13	data.truthValue.....	44
B.4.14	data.word	45
B.4.15	data.xTruthValue.....	45
Annex C (normative) Morpho-syntactic data categories		46
Bibliography.....		58

Foreword

ISO (the International Organization for Standardization) is a worldwide federation of national standards bodies (ISO member bodies). The work of preparing International Standards is normally carried out through ISO technical committees. Each member body interested in a subject for which a technical committee has been established has the right to be represented on that committee. International organizations, governmental and non-governmental, in liaison with ISO, also take part in the work. ISO collaborates closely with the International Electrotechnical Commission (IEC) on all matters of electrotechnical standardization.

International Standards are drafted in accordance with the rules given in the ISO/IEC Directives, Part 2.

The main task of technical committees is to prepare International Standards. Draft International Standards adopted by the technical committees are circulated to the member bodies for voting. Publication as an International Standard requires approval by at least 75 % of the member bodies casting a vote.

Attention is drawn to the possibility that some of the elements of this document may be the subject of patent rights. ISO shall not be held responsible for identifying any or all such patent rights.

ISO 24611 was prepared by Technical Committee ISO/TC 37, *Terminology and other language and content resources*, Subcommittee SC 4, *Language resource management*.

iteh Standards
(<https://standards.iteh.ai>)
Document Preview

ISO 24611:2012

<https://standards.iteh.ai/catalog/standards/iso/442a57b7-65de-4d5f-9c8f-34c401137730/iso-24611-2012>

Introduction

ISO/TC 37/SC 4 focuses on the definition of models and formats for the representation of annotated language resources. To this end, it has generalised the modelling strategy initiated by its sister committee, SC 3, for the representation of terminological data [Romary, 2001], through which linguistic data models are seen as the combination of a generic data pattern (a meta-model), which is further refined through a selection of data categories that provide the descriptors for this specific annotation level. Such models are defined independently of any specific formats, and ensure that an implementer has the necessary conceptual instrument with which to design and compare formats with regard to their degrees of interoperability.

One important aspect of representing any kind of annotation is the capacity to provide a clear and reliable semantics for the various descriptors used, either in the form of formal features and feature values, or directly as objects in a representation that is expressed, for instance, in XML. In order to be shared across various annotation schemas and encoding applications, such a semantics should be implemented as a centralised registry of concepts: we will henceforth refer to these as data categories. As such, data categories should bear the following constraints.

- From a technical point of view, they must provide unique, stable references (implemented as persistent identifiers, in the sense of ISO 24619) such that the designer of a specific encoding schema can refer to them in his or her specification. By doing so, two annotations will be deemed to be equivalent when they are in fact defined in relation to the same data categories (as feature and feature value).
- From a descriptive point of view, each unique semantic reference should be associated with precise documentation combining a full text elicitation of the meaning of the descriptor with the expression of specific constraints that bear upon the category.

In recent years, ISO has developed a general framework for representing and maintaining such a registry of data categories, encompassing all domains of language resources. This initiative, described in ISO 12620, has led to the implementation of an online environment providing access to all data categories that have been standardized in the context of the various language resource-related activities within ISO, or specifically as part of the maintenance of the data category registry. It also provides access to the various data categories that individual language technology practitioners have defined in the course of their own work and decided to share with the community.

The ISO data category registry, as available through the ISOCat (www.isocat.org) implementation, is intended as a 'flat' marketplace of semantic objects, providing only a limited set of ontological constraints. The objective there is to facilitate the maintenance of a comprehensive descriptive environment where new categories are easily inserted and reused without the need for any strong consistency check with the registry at large. Indeed, the following basic constraints are part of the data category model, as defined in ISO 12620:

- simple generic-specific relations, when these are useful for the proper identification of interoperability descriptors between data categories. For instance, the fact that /properNoun/ is a sub-category of /noun/ makes it possible to compare morpho-syntactic annotations based on different descriptive levels of granularity;
- the description of conceptual domains, in the sense of ISO 11179, to identify, when known or applicable, the possible value of so-called complex data categories. For instance, it can be used to record that possible values of /grammaticalGender/ (limited to a small group of languages [Romary 2011]), could be a subset of {/masculine/, /feminine/ and /neutral/};
- language-specific constraints, either in the form of specific application notes or as explicit restrictions bearing upon the conceptual domains of complex data categories. For instance, it is possible to express explicitly that /grammaticalGender/ in French can only take the two values: {/masculine/ and /feminine/}.

This International Standard provides a comprehensive framework for the representation of morpho-syntactic (also referred to as part-of-speech) annotations. Such an annotation level corresponds to a first lexical abstraction level over language data (textual or spoken) and, depending on the language to be annotated, together with the characteristics of the annotation tool or annotation scheme that is being used, can vary enormously in structure and complexity.

In order to deal with such complex issues as ambiguity and determinism in morpho-syntactic annotation, this International Standard introduces a meta-model that draws a clear distinction between the two levels of tokens (representing the surface segmentation of the source) and word-forms (identifying lexical abstractions associated with groups of tokens). These two levels share the following specificities: on the one hand, they can be represented as simple sequences and as local graphs such as multiple segmentations and ambiguous compounds; on the other hand, any n-to-n combination can stand between word forms and tokens.

As linguistic segments (sometimes called ‘markables’ in the literature [see, for instance, Carletta et al. 1997]), *tokens* may be embedded in the source document as inline mark-up, or they may point remotely to it by means of so-called stand-off annotations.

As linguistic abstractions, *word-forms* can be qualified by various linguistic features characterising the morpho-syntactic properties that are instantiated in the realisation of the lexical entry within the annotated text. Such properties may range from the simple indication of a lemma up to an explicit reference to a lexical entry in a dictionary. In most existing applications of morpho-syntactic annotation, linguistic properties are expressed by means of so-called tags; these codes refer to basic feature structures (see early examples in Monachini and Calzolari, 1994). Such codes may also provide morphological information, including its part of speech (e.g. noun, adjective or verb), and features such as number, gender, person, mood and verbal tense.

In keeping with the general modelling strategy of ISO/TC 37, this International Standard/MAF provides means of relating morpho-syntactic tags expressed as feature structures (compliant with ISO 24610) to the data categories available in ISOCat. A normative annex of this International Standard elicits a core set of data categories that can be used as reference for most current morpho-syntactic annotation tasks in a multilingual context. However, when implementers of this International Standard find these categories inappropriate in either coverage, scope or semantics, they are encouraged to use ISOCat to define their own categories in compliance with ISO/TC 37 principles.

Associated to the meta-model, MAF also provides a default XML syntax that may be used to serialise MAF-compliant annotation models. Since many existing projects are based on the text encoding initiative (TEI) guidelines (www.tei-c.org) — particularly in digital humanities, where a proper encoding of textual sources is essential — this International Standard will also provide clues about how to articulate the MAF model with TEI-compliant encodings. Indeed, the TEI guidelines already offer a variety of constructs and mechanisms to cope with many issues relevant to spoken corpora and their annotations (Romary and Witt, 2012).

Finally, it should be noted here that this International Standard forms the conceptual basis for the development of the ISO 24614 series on word segmentation, whereby all general principles and rules defined in ISO 24614-1, as well as the constraints expressed in additional parts for specific languages, are to be understood according to the token–word-form dichotomy.

Language resource management — Morpho-syntactic annotation framework (MAF)

1 Scope

This International Standard provides a framework for the representation of annotations of word-forms in texts; such annotations concern tokens, their relationship with lexical units, and their morpho-syntactic properties.

It describes a metamodel for morpho-syntactic annotation that relates to a reference to the data categories contained in the ISOCat data category registry (DCR, as defined in ISO 12620). It also describes an XML serialization for morpho-syntactic annotations, with equivalences to the guidelines of the TEI (text encoding initiative).

2 Normative references

The following referenced documents are indispensable for the application of this document. For dated references, only the edition cited applies. For undated references, the latest edition of the referenced document (including any amendments) applies.

ISO 24610-1, *Language resource management — Feature structures — Part 1: Feature structure representation*

3 Terms and definitions

For the purposes of this document, the terms and definitions given in ISO 24610-1 and the following apply.

3.1

DAG

directed acyclic graph

graph with directed edges and no cycles

Note 1 to entry: DAGs are a subset of *finite state automata* (3.4).

3.3

feature structure

set of feature specifications, used in the morpho-syntactic annotation framework (MAF) to express morpho-syntactic content

Note 1 to entry: Feature structures are described in ISO 24610-1.

3.4

FSA

finite state automata

graphs made up of states with an initial state and a final state, and a finite set of transitions from state to state

Note 1 to entry: See also *DAG* (3.1).

3.5

grapheme

minimal unit in a written language

EXAMPLE Letter, pictogram, ideogram, numeral, punctuation.

3.6

inflection

modification or marking of a lexeme that reflects its morpho-syntactic properties

3.7

inflected form

form that a word can take when used in a sentence or a phrase

Note 1 to entry: An inflected form of a word is associated with a combination of morphological features, such as grammatical number and case.

3.8

lemma

lemmatised form

conventional form chosen to represent a lexeme

Note 1 to entry: In European languages, the lemma is usually the *singular* if there is a variation in number, the *masculine* form if there is a variation in *gender*, and the *infinitive* for all verbs. In some languages, certain nouns are defective in the singular form; in these cases, the plural is chosen. For verbs in Arabic, the lemma is usually deemed to be the third person singular with the accomplished aspect.

3.9

lexeme

morpheme generally associated with a set of word-forms sharing a common meaning

3.10

lexical entry

container for managing a set of word-forms and possibly one or more meanings to describe a lexeme

<https://standards.iteh.ai/catalog/standards/iso/442a57b7-65de-4d5f-9c8f-34c401137730/iso-24611-2012>

lexicon

resource comprising a collection of lexical entries for a language

3.12

morpheme

smallest linguistic unit that carries a meaning in a discourse, but which cannot be divided into smaller meaningful units

Note 1 to entry: A morpheme is either grammatical (grammeme) or lexical (lexeme).

3.13

morphological feature

morpho-syntactic feature

feature induced from the inflected form of a word

Note 1 to entry: The ISOCat data category registry provides a comprehensive list of values for European languages.

EXAMPLE “grammaticalGender”.

3.14

morphology

description of the structure and formation of word-forms

3.15**morpho-syntactic tag****tag**

feature structure used systematically to qualify a word-form

3.16**tagset**

comprehensive set of tags used for the morpho-syntactic description of a language

Note 1 to entry: The ISOCat data category registry is to be used as the reference for describing a tagset.

3.17**part of speech****grammatical category**

category assigned to a word based on its grammatical and semantic properties

EXAMPLE Noun, verb.

Note 1 to entry: The ISOCat data category registry provides a comprehensive list of values for parts of speech.

3.18**phoneme**

minimal unit in the sound system of a language

3.19**script**

set of graphic characters used for the written form of one or more languages

3.20**syntagmatic relation**

relation by which linguistic units in a discourse are associated

3.21**token**

non-empty contiguous sequence of graphemes or phonemes in a document

Note 1 to entry: For editorial reasons, some annotation scheme may extend the notion of token to an empty sequence. See the section on token attachment (6.2).

3.22**tokenization**

process identifying tokens

3.23**transcription**

form resulting from a coherent method of writing down speech sounds

3.24**transliteration**

form resulting from the conversion of one script into another, usually through a one-to-one correspondence between characters

3.25**word-form****morpho-syntactic unit**

contiguous or non-contiguous linguistic unit identified as corresponding to a lexical entity in a language

Note 1 to entry: Word-forms may have no acoustic or graphic realization, or may correspond to one or more tokens.

3.26

word lattice

set of possible alternative decompositions of a text or speech segment into word-forms

Note 1 to entry: A word lattice has the algebraic properties of a directed acyclic graph with an initial node and a final node.

Note 2 to entry: See also *DAG* (3.1) and *FSA* (3.4).

4 The MAF meta-model

4.1 Overview

Morpho-syntactic annotations provide an important layer of linguistic information in a document. This International Standard is based on a meta-model that draws a clear distinction between the two levels of tokens (representing the surface segmentation of the source) and word-forms (identifying lexical abstractions associated to groups of tokens). These two levels share the following specificities: on the one hand, they can be represented as simple sequences and as local graphs (e.g. multiple segmentations and ambiguous compounds); on the other hand, any n-to-n combination can stand between word-forms and tokens. This International Standard delimits minimal and maximal sequences in documents (either text or speech) that can be identified as word-forms and seeks to categorise the linguistic and distributional criteria that may be used to mark these word-forms within some larger syntagmatic context. Minimal units cannot be further decomposed using similar criteria, but may however be divided into smaller units using morphological or phonological properties. Word-forms can be aggregated to form maximal units (such as compound words or multi-word units) that act as elementary units for other levels of linguistic analysis, particularly syntax. In particular, word-forms correspond to the non-terminal level defined in ISO 24615.

4.2 MAF Meta-model

Figure 1 presents a simplified view of the proposed meta-model for morpho-syntactic annotations, whereas Figure 2 presents a more formal view based on UML (Unified Modeling Language).

[ISO 24611:2012](https://standards.iteh.ai/ISO/24611/2012)

<https://standards.iteh.ai/catalog/standards/iso/442a57b7-65de-4d5f-9c8f-34c401137730/iso-24611-2012>

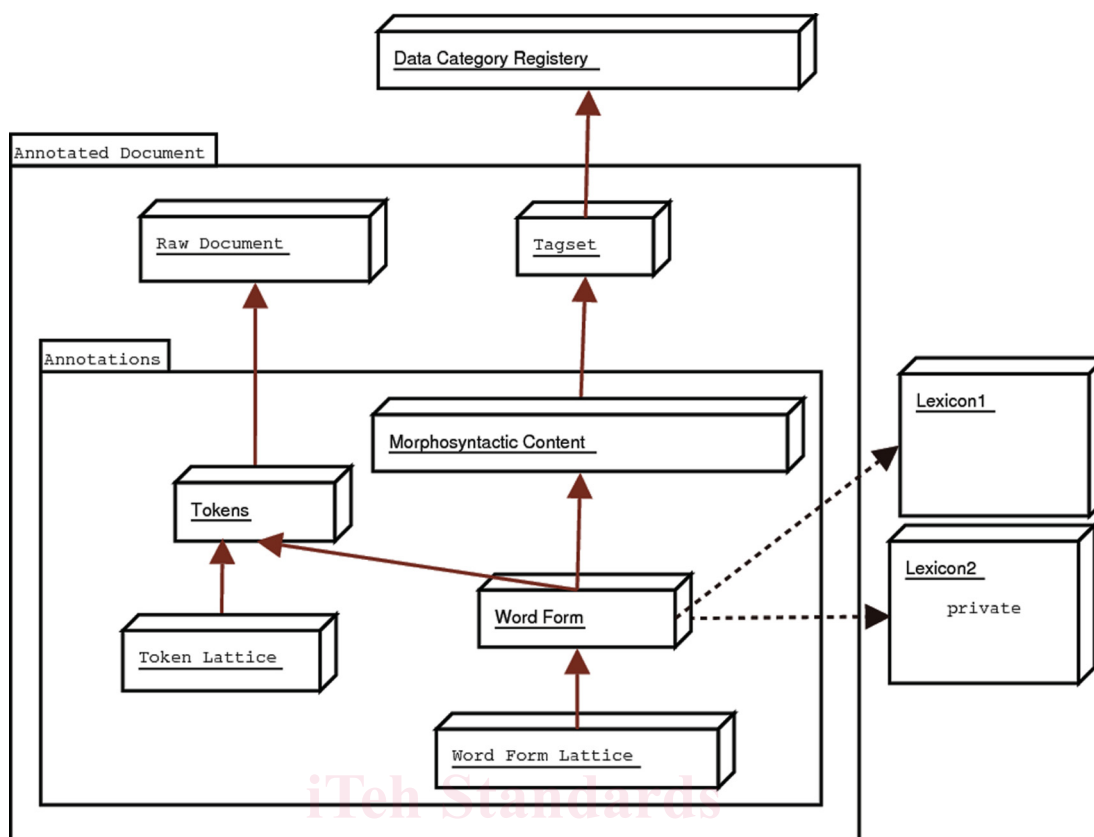


Figure 1 — Simplified view of MAF meta-model

An annotated document comprises an original document and a set of annotations. Annotations are associated with word-forms corresponding to zero or more tokens in the original document. A word-form may also be associated with a lexical entry providing information about its underlying lemma and inflected form. The morpho-syntactic annotation associated with a word-form is represented by a tag, the significance of which may be expressed as a feature structure. A set of such tags used by a particular annotation scheme is referred to as a tagset, and corresponds to what is defined in the ISO 24610-2-specified feature structures representation (FSR) as a feature structure library. Each discrete category within such a tagset should be describable in terms of registered data categories as described in ISO 12620 and implemented in ISOCat. Because annotation may be applied both to tokens and to word-forms, structural ambiguity is likely. Hence annotation is typically conceptualised as one or more streams, each represented as a word lattice or more formally as a directed acyclic graph (DAG).

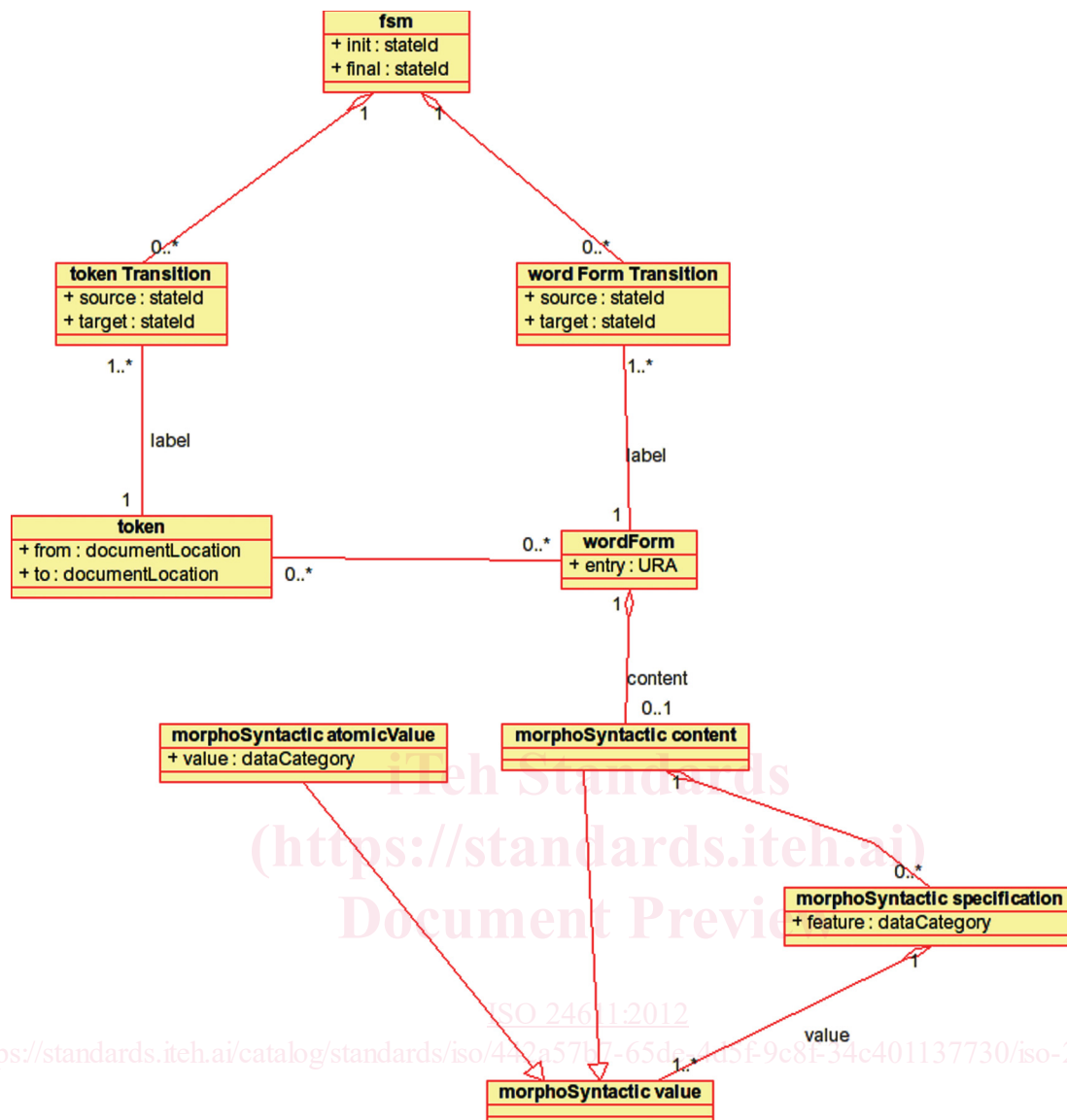


Figure 2 — UML view of MAF meta-model

5 Segmenting with tokens

5.1 General

Morpho-syntactic annotations are carried by segments, called tokens, that are present in the document flow, but this does not imply that the resulting segmentation corresponds to a sequence of adjacent segments partitioning the original document. It is particularly important to distinguish the word-forms from their realisations. Some parts of a document may carry no annotations (e.g. typographic marks, stage directions and markup elements) while other parts may not correspond exactly to their segmented form (e.g. abbreviations, brachygraphies, orthographic errors and variations, and typographic and morphological contractions). A word-form may not correspond exactly to a segment identified by orthographic marks such as white spaces or hyphens (e.g. for German compound words, speech transcription and Sanskrit writing).