

Classification of German job titles in online job postings using the KldB- 2010 taxonomy Model-Version 3

Technical Report
Bertelsmann Stiftung

Version

1

Authors

Johannes Müller

Erstellt am

29. Juni 2026

Table of Contents

1	Executive Summary	3
2	Task Definition	4
3	Data & Annotation.....	5
4	Taxonomy Representation	9
5	Model & Inference Pipeline	10
6	Training	12
7	Evaluation	14
8	Production Validation.....	17
9	Limitations, Risks & Recommendations	20

Executive Summary

This method report describes the dataset, annotation process, taxonomy representation, model pipeline, evaluation, and production validation for mapping German online job vacancy titles to KLDB-2010 occupation codes.

Summary: The Occupation Classifier v3 is a German job-title to KLDB-2010 occupation classifier. It is framed as a retrieval task: a fine-tuned bi-encoder embeds a job title and retrieves the closest of ~4,400 KLDB 8-Steller occupation codes from a structured FAISS taxonomy index. Retrieval runs at the 8-Steller level, while the primary evaluation metric is at the 5-Steller level (Berufsgattung).

Dataset: The model is trained and evaluated on ~13,300 usable gold-standard (job title, KLDB code) pairs, human-annotated from a sample of 400M+ German online vacancy titles from TextKernel, 2022–2026. The source data combines three samples: General (~76%), Green Jobs (~13%), and AI/KI (~12%). The dataset was built through an iterative annotate → train → cleanup loop, including adjudication of 982 model/annotation disagreements, of which 603 gold labels were corrected.

Model: The classifier uses a bi-encoder fine-tuned from mxbai-embed-de-large-v1 (560M parameters, 1,024-dimensional embeddings, cosine similarity). The training setup uses MultipleNegativesRankingLoss with hard negatives. The production pipeline adds a rule-based pre-filter for unclassifiable titles and a 0.4 cosine-similarity threshold below which predictions return null. Each KLDB 8-Steller code is represented by one structured taxonomy entry (Titel: ... / Ähnlich: ...) combining the canonical occupation title with its unique synonyms. This representation yielded **+12.2pp P@1** over a multi-anchor baseline and is the single biggest improvement in the project.

Results: The final model reaches **P@1 kldb5 = 0.847** and **R@5 kldb5 = 0.955**. This is **+37.1pp** over the un-tuned base model, with hard negatives adding **+4.7pp** on top of gold-only fine-tuning.

Production behavior: On a 100k-title production sample, ~97.5% of titles are kept after filtering and thresholding. Spot-checks of the 20 most frequent job titles show plausible predictions, and targeted difficult-code checks show that the new model corrects known legacy-model error patterns.

Limitations: The model is German/KLDB-2010 only. It is not directly transferable to ISCO, SOC, or other national classification systems. Known limitations include taxonomy coverage, long-tail class imbalance, temporal and platform bias, generic-title ambiguity, and score-threshold calibration.

1 Task Definition

1.1 Classification Goal

The goal is to classify short German job vacancy titles into standardized KLDB-2010 occupation codes. The input is a job title, for example “Softwareentwickler/in (m/w/d)”. The output is a KLDB occupation concept, represented by a KLDB 8-Steller code and evaluated primarily at the KLDB 5-Steller level.

1.2 Retrieval Formulation

The task is formulated as a dense retrieval problem:

1. Encode the job title as a query embedding.
2. Search a FAISS index of structured KLDB taxonomy entries.
3. Return the nearest occupation-code entry.
4. Aggregate or evaluate the result at the relevant KLDB level.

Retrieval is performed at the 8-Steller level because this is the most specific target available in the taxonomy index. The primary evaluation is performed at the 5-Steller level because this corresponds to the Berufsgattung level and is the main target for downstream use.

1.3 KLDB Levels and Evaluation Target

The report uses several KLDB aggregation levels:

Level	Description	Use in this report
KLDB-8 / BKZ	Berufsbezeichnung (Most specific occupation code)	Retrieval target
KLDB-5	Berufsgattung	Primary evaluation target
KLDB-3	Berufsgruppe (Broader occupation group)	Secondary metric, especially useful for internships and student roles

2 Data & Annotation

2.1 Source Data

The job titles are sampled from a TextKernel dataset of over 400 million German online job vacancy titles, collected as part of the Bertelsmann Stiftung DOSTA project. Three targeted samples were drawn to ensure broad coverage across general vacancies and strategically important domains:

Sample	Size	Selection method
General	~10,000	Random sample from the full dataset, with 2x oversampling of non-TOP-100 KLDB codes to increase representation of rare occupations
Green Jobs	2,000	Job titles pre-classified as sustainability-related by a separate model
AI / KI	~2,000	Job titles pre-classified as AI-related by a separate model

The dataset contains only job titles and occupation codes. Job posting UUIDs in the evaluation set are anonymized identifiers. Occasionally, a company name may occur in a job title.

2.2 Dataset Composition

The cleaned dataset supports three core artifacts:

Dataset	Records	Description
Training	20,792	Training triplets: query, positive, hard negative
Evaluation	2,468	Evaluation set with gold-standard KLDB labels
Taxonomy	4,417	Structured taxonomy of all KLDB 8-Steller codes

After excluding unusable, ambiguous, and placeholder records, the final supervised dataset contains 13,303 usable pairs, split into 10,325 train and 2,468 eval examples. The split is stratified by KLDB 5-Steller frequency; codes with only one example go to training.

2.3 Annotation Process

A gold-standard dataset of ~14,500 (job_title, bkz_id) pairs was created by presenting human annotators with candidate labels from both the legacy DOSTA model and an untrained base embedding model.

For each title, annotators selected one of up to four candidate codes or marked the title as unusable.

Annotation value	Meaning
1-4	Annotator selected one of the provided BKZ options
9	Annotators disagreed; resolved via separate adjudication sheet
99	Title marked as unusable or unclassifiable

Annotations were performed by domain experts at &effect data solutions GmbH with knowledge of the KLDB-2010 classification system.

2.4 Cleanup and Adjudication

The dataset was created through an iterative annotation-training-cleanup pipeline:

1. [Annotation from scratch](#) — Human annotators reviewed every case, resolved disagreements via adjudication, and marked ambiguous or unusable titles.
2. [Initial model training](#) — A sentence embedding model was fine-tuned on gold pairs against the structured taxonomy index.
3. [Hard-negative mining](#) — The model was used to generate hard negatives for a second training round.
4. [Annotation cleanup](#) — The best model was run against the full annotated dataset. All model/annotation disagreements were reviewed.
5. [Retrain on clean data](#) — After cleanup, the full pipeline was re-run from scratch.

The cleanup step reviewed [982 disagreements](#):

Outcome	Count	Interpretation
Gold label corrected	603	Model prediction was judged correct; the original annotation was wrong
Ambiguous	47	Title was genuinely ambiguous
Original label confirmed	332	The original annotation was retained

2.5 Special-Case Flagging

Titles containing keywords for internships, working students, student theses, or dual study programmes were flagged with `is_special_case = true`.

Keyword	Category
praktikum, praktikant, praktikantin, internship	Internship
werkstudent, werkstudentin	Working student
masterarbeit, bachelorarbeit, abschlussarbeit, diplomarbeit, studienarbeit, thesis, phd	Student thesis / research
duales studium	Dual study programme

These titles are not necessarily unusable. They are flagged because their occupational target is often broader or less stable than for regular job titles.

2.6 Annotation Rules for Ambiguous Titles

Domain experts established explicit mapping rules for frequently occurring ambiguous or generic job titles.

The guiding principles were:

- Generic titles → most likely occupation. When a job title is ambiguous, the most common real-world occupation in that category is chosen.
- Context decides. The same title may receive different codes depending on domain context.
- Too generic → code 99 or placeholder. When no sensible default exists, the title is marked as unusable or assigned a placeholder code.
- Helper roles map to the most likely helper occupation.
- Custom placeholder codes are used only where the standard KLDB taxonomy has no appropriate generic code.

The default mappings for generic titles were calibrated against the BA Jobbörse in March 2026, where the full posting context is available.

2.6.1 Fixed mappings for common generic titles

Generic title	KLDB code	Mapped occupation
Analyst (unspecified)	7131410400	Business-Analyst/in
Assistenz / Büroassistentz	7140211500	Büroassistent/in
Bauhelfer / Baustellenmitarbeiter	3210110000	Helfer/in - Hochbau
Consultant (generic / IT)	4322312400	IT-Berater/in

Generic title	KLDB code	Mapped occupation
Disponent (unspecified)	5162210100	Speditionskaufmann/-frau
Elektriker	2621210300	Elektroinstallateur/in
Elektroniker (generic)	2625218900	Elektroniker/in - Betriebstechnik
Mechaniker / Monteur (generic)	3434210300	Montagemechaniker/in
Sachbearbeiter (unspecified)	7140210100	Bürokaufmann/-frau
Servicetechniker / Techniker (generic)	2513213200	Kundendienstmonteur/in
Researcher / PostDoc / PhD	8430419000	Wissenschaftliche/r Mitarbeiter/in
Software Engineer / Developer	4341410200	Softwareentwickler/in

2.6.2 Context-dependent mappings

Title	Context	KLDB code	Mapped occupation
Architekt	Construction	3111410300	Architekt/in
Architekt	IT / enterprise / solution architecture	4341310300	IT-Lösungsentwickler/in
Consultant	Tax	7230410000	Steuerberater
Consultant	Automotive	6128410100	Betriebswirt/in - Automobilwirtschaft
Disponent	Route planning / freight	5163210700	Disponent/in - Güterverkehr
Fahrer	LKW / long-haul	5212210800	Berufskraftfahrer/in
Fahrer	Passenger transport	5211211000	Fahrer/in - Personenbeförderung
Fahrer	Delivery / courier	5218210100	Auslieferungsfahrer/in

Examples of unresolvable titles marked as code 99 include: Entwickler, Helfer (allgemein), Ingenieur, Manager, Naturwissenschaftler, Polier, and Vorarbeiter without qualifying context.

2.7 Source Data Quality

Dimension	Facts	Implication
Quality	92% clean; 3.6% unusable; 4.3% ambiguous	High signal-to-noise for a human annotation task on short text
Coverage	1,761 of 4,415 BKZ codes covered by examples	~40% of the taxonomy has annotated examples
Class distribution	Median 3 examples per class; max 200	Extreme long tail
Special cases	461 internship/thesis titles flagged	Can be filtered or separately evaluated for regular-job use cases
Placeholders	106 generic codes such as 20001xx and XXXX3/4	Excluded from standard evaluation

3 Taxonomy Representation

3.1 KLDB-2010 Source

The taxonomy is based on the Bundesagentur für Arbeit (BA) classification of occupations, KLDB-2010, including official searchwords and synonyms.

The retrieval index contains all currently used BKZ codes and their occupation titles. Since some searchwords are ambiguous and appear under multiple occupation codes, only unique searchwords assigned to exactly one BKZ code are used in the structured synonym list.

3.2 Structured Taxonomy Entries

Instead of using ~102k short individual anchor strings, one per searchword, each occupation code is represented by one rich structured text entry:

Titel: Softwareentwickler/in
Ähnlich: Programmierer, Anwendungsentwickler, Backend-Entwickler, ...

This design gives the model more discriminative context per code by combining the canonical title and its unique synonyms in a single passage.

3.3 Comparison to Multi-Anchor Baseline

The structured representation yielded [+12.2pp P@1](#) over the multi-anchor baseline. This makes taxonomy representation the single largest methodological improvement in the project.

4 Model & Inference Pipeline

4.1 Base Model

The classifier is fine-tuned from mixedbread-ai/deepset-mxbai-embed-de-large-v1.

Property	Value
Model type	Bi-encoder sentence embedding model
Language	German, with some English job titles
Parameters	560M
Output dimensionality	1,024
Similarity function	Cosine similarity
Maximum sequence length	512 tokens

The model uses query: and passage: prefixes inherited from the mxbai base model.

4.2 Inference Pipeline

The production inference pipeline is:

1. Apply a conservative rule-based pre-filter.
2. Encode the job title as a query embedding.
3. Search the FAISS index of structured KLDB taxonomy entries.
4. Return the nearest KLDB concept if the similarity score is at least 0.4.
5. Return concept: null with a filter_reason if the title is filtered or below threshold.

4.3 Rule-Based Pre-Filter

The filter removes clearly unclassifiable inputs before embedding search. It is intentionally conservative: borderline-ambiguous titles are passed to the model and may be caught by the score threshold instead.

Layer	Filter	Reason code	Examples
Structural	Encoding artifacts, non-Latin script, too short, date strings, truncated fragments, short English fragments	encoding_error, non_german_script, too_short, date_string, truncated_fragment, english_fragment	BÃ¼rokaufmann, Arabic/Cyrillic titles, m/w/d, 2024-01-15
Language	Polish job keywords, gender marker only	non_german_language, gender_marker_only	praca spawacz niemcy, (m/w/d)
Ausbildung	Bare "Ausbildung" or "Ausbildung YEAR" without a specific trade	ausbildung_year	Ausbildung 2026 (m/w/d), Schule & Ausbildung
Internship	Bare student or intern roles without a field	internship_or_student, schueler_praktikum	Werkstudent (m/w/d), Praktikum, Schülerpraktikum Büro
Other roles	Bare FSJ/BFD, Minijob/Nebenjob, Initiativbewerbung	volunteer_service, minijob_nebenjob, initiativbewerbung	FSJ (m/w/d), Minijob (m/w/d), Initiativbewerbung
Content	Marketing/spam sentences, excessive punctuation, UI or scraping artifacts	marketing_sentence, spam, ui_artifact	Kommen Sie zu uns!, JETZT BEWERBEN!!!, Navigation überspringen
Generic roles	Bare generic job titles without qualifying context	generic_*	Mitarbeiter/in, Helfer (m/w/d), Aushilfe, Trainee

Titles with a qualifying field after a generic role are **not** filtered. Examples that pass through include Sachbearbeiter Buchhaltung, Werkstudent Online Marketing, and Ausbildung 2026 zum Bäcker.

4.4 Score Threshold and Null Predictions

After FAISS retrieval, a minimum cosine-similarity cutoff of **0.4** is applied. Predictions below this threshold are returned as concept: null with filter_reason: "low_confidence".

Calibration on 100k production titles from 2025-01 to 2026-02:

Outcome	Count	Share
Already filtered by rules	650	0.7%
Newly filtered by threshold (score < 0.4)	1,803	1.8%
Kept (score $\geq$ 0.4)	97,547	97.5%

The low-confidence titles were almost exclusively garbage, such as bare years, UI fragments, non-German text, or company names.

Confidence interpretation:

Score range	Interpretation
> 0.85	High confidence; likely correct
0.4–0.85	Medium confidence; useful for monitoring and spot-checking
< 0.4	Returned as null

5 Training

5.1 Training Data

Training uses [20,792 triplets](#):

- 10,396 gold query-positive pairs
- 10,396 hard-negative triplets

Each training example contains:

1. A job title as the query.
2. The correct structured taxonomy entry as the positive passage.
3. A hard-negative taxonomy entry retrieved by an earlier model but judged incorrect.

5.2 Training Objective

The model is fine-tuned using MultipleNegativesRankingLoss with cosine similarity.

Parameter	Value
Loss function	MultipleNegativesRankingLoss, scale=20.0
Epochs	4, best at epoch 4

Parameter	Value
Batch size	128
Learning rate	2e-5
LR scheduler	Linear with warmup
Warmup ratio	0.1
Optimizer	AdamW, fused
Weight decay	0.0
Max gradient norm	1.0
Training regime	bf16 mixed precision
Gradient checkpointing	Enabled
Seed	42

5.3 Token Statistics

Column	Min tokens	Mean tokens	Max tokens
Anchor / job title	4	16.4	59
Positive / taxonomy entry	9	233.7	512

5.4 Training Loss

Epoch	Step	Training loss
0.61	100	1.3062
1.23	200	0.5745
1.84	300	0.4640
2.45	400	0.3824
3.07	500	0.3557
3.68	600	0.3261

5.5 Compute Environment

Component	Value
Hardware	NVIDIA A100 40GB via Modal
Python	3.12.1
sentence-transformers	5.2.3
Transformers	4.57.6
PyTorch	2.10.0+cu128
Accelerate	1.13.0
Datasets	4.7.0
Tokenizers	0.22.2
Vector search	FAISS, faiss-cpu

6 Evaluation

6.1 Evaluation Split

The evaluation set contains [2,484 examples](#) with gold-standard KLDB labels. It is drawn from all three source samples:

Sample	Eval rows	Share
General	1,883	76%
Green Jobs	313	13%
AI / KI	288	12%
Total	2,484	

6.2 Metrics

Metric	Definition
P@1 kldb5	Top-1 prediction matches at 5-Steller level; primary metric

Metric	Definition
P@1 kldb8	Top-1 prediction matches at 8-Steller / BKZ level
P@1 kldb3	Top-1 prediction matches at 3-Steller level
R@5 kldb5	Correct 5-Steller appears anywhere in the top 5
R@5 kldb8	Correct BKZ appears anywhere in the top 5
R@5 kldb3	Correct 3-Steller appears anywhere in the top 5

6.3 Overall Results

Model	P@1 kldb5	P@1 kldb8	P@1 kldb3	R@5 kldb5	R@5 kldb8	R@5 kldb3
mxbai-de-large, base, no fine-tuning	0.476	0.358	0.608	0.717	0.609	0.839
mxbai-de-large, fine- tuned, gold-only	0.800	0.729	0.849	0.946	0.927	0.957
mxbai-de-large, fine- tuned, gold + hard negatives	0.847	0.782	0.881	0.955	0.939	0.963

Fine-tuning improves P@1 kldb5 by **+37.1pp** over the base model. Hard-negative mining adds **+4.7pp** on top of gold-only fine-tuning.

6.4 Per-Epoch Progression

Epoch	P@1 kldb5, gold-only	P@1 kldb5, gold + hard negatives
1	0.773	0.830
2	0.782	0.841
3	0.798	0.846
4	0.800	0.847

6.5 Performance by Training Frequency

Bucket	Description	N	P@1 kldb5	P@1 kldb8	R@5 kldb5	R@5 kldb8
B	1–3 training examples	195	0.728	0.692	0.903	0.892
C	4–10 training examples	210	0.819	0.743	0.943	0.929
D	10+ training examples	2,079	0.861	0.795	0.961	0.944
Overall		2,484	0.847	0.782	0.955	0.939

Performance improves with more training examples but remains strong even for codes with only 1–3 examples.

6.6 Performance by Sample Type

Sample	N	P@1 kldb5	P@1 kldb3	R@5 kldb5
General	1,883	0.843	0.880	0.955
Green Jobs	313	0.916	0.926	0.974
AI / KI	288	0.795	0.840	0.927

Green Jobs titles are easiest to classify. AI/KI titles are the most challenging sample.

6.7 Special-Case Titles

A separate evaluation was conducted on 460 titles flagged as special cases, such as Praktikum Data Science or Werkstudent Marketing.

For these titles, KLDB-3 is especially relevant because the broad occupational group often matters more than the exact 8-Steller code.

Metric	Value
P@1 kldb3	0.769
P@1 kldb5	0.562
R@5 kldb3	0.935
R@5 kldb5	0.845

7 Production Validation

7.1 100k Production Sample

The model was run on 100,000 job titles from the production database, sampled from 2025-01-01 to 2026-02-28.

Outcome	Count	Share
Total titles	100,000	
Filtered by pre-processing	650	0.7%
Classifiable	99,350	99.4%
Same as old model	58,428	58.8% of classifiable
Changed	40,922	41.2% of classifiable

The 41% prediction-change rate is expected because the new model is trained on different data and uses a different index structure.

7.2 Distribution Comparison

In the 100k production sample, several occupations that had previously been under-classified by the legacy model enter the top 20, including:

- IT-Kundenbetreuer/in (43223)
- Vertriebsassistent/in (61122)
- Elektroniker/in - Betriebstechnik (26252)
- Systemprogrammierer/in (43413)

The top-ranked new-model categories are:

#	KLDB-5	Count	Change vs. old model	Occupation
1	62102	3,259	=	Kaufmann/-frau - Einzelhandel
2	71402	2,126	=	Büro- & Sekretariatskräfte
3	51311	1,666	=	Lagermitarbeiter/in
4	72213	1,566	+3	Buchhalter/in
5	51312	1,511	=	Kommissionierer/in
6	43223	1,425	NEW	IT-Kundenbetreuer/in

#	KLDB-5	Count	Change vs. old model	Occupation
7	54101	1,333	+3	Helfer/in - Reinigung
8	52122	1,163	+1	Berufskraftfahrer/in
9	82102	1,137	+2	Altenpfleger/in
10	61122	1,132	NEW	Vertriebsassistent/in

7.3 Frequent-Title Spot Check

Predictions for the 20 most frequent job titles in production were manually reviewed. All predictions are plausible.

Representative examples:

Job title	KLDB-5	Predicted occupation	Score
Produktionsmitarbeiter (m/w/d)	20001	Produktionsmitarbeiter, custom code	0.882
Reinigungskraft (m/w/d)	54101	Reinigung	0.849
Verkäufer (m/w/d)	62102	Verkauf	0.896
Staplerfahrer (m/w/d)	52531	Kranführer/Aufzugsmaschinisten	0.880
Ausbildung 2026	—	filtered: ausbildung_year	—
Elektriker (m/w/d)	26212	Bauelektrik	0.786
Elektroniker (m/w/d)	26252	Elektrische Betriebstechnik	0.787

The Elektriker and Elektroniker mappings match the annotation rules.

7.4 Difficult-Code Analysis

The new model was tested on 1,000 job titles previously classified by the legacy model into five known problematic KLDB-5 codes.

Legacy KLDB-5	Occupation	Same	Changed	Filtered
26243	Regenerative Energietechnik	9 (4%)	190 (95%)	1
31114	Architekt:innen	129 (65%)	70 (35%)	1
61124	Vertriebler:innen, außer IKT	124 (62%)	76 (38%)	0

Legacy KLDB-5	Occupation	Same	Changed	Filtered
84304	Hochschullehre	93 (47%)	103 (52%)	4
92113	Werbung / Marketing	98 (49%)	102 (51%)	0

Key findings:

- **26243 Regenerative Energietechnik:** 95% reclassified. The legacy model had pulled in generic Elektriker and Servicetechniker titles. The new model correctly routes many of them to more appropriate categories.
- **31114 Architekt:innen:** 65% stay. The changed cases are mostly IT architects, now routed to software or IT solution development.
- **61124 Vertrieb hoch komplex:** 62% stay. Reclassified titles move mostly to lower complexity levels within Vertrieb.
- **84304 Hochschullehre:** 47% stay. Teaching and training titles are redistributed to more specific education codes.
- **92113 Werbung / Marketing:** About half stay; the rest move across marketing leadership and other marketing subcategories.

7.5 Regional Analysis

The model was run on **15,000 job titles**, 5,000 each from Hannover, München, and Bielefeld.

Filter rates:

Region	Titles	Filtered	Rate
Hannover	5,000	125	2.5%
München	5,000	166	3.3%
Bielefeld	5,000	126	2.5%

The regional differences in the top categories are consistent with the economic profiles of the regions: München shows higher shares in IT, software, marketing, and white-collar occupations; Hannover and Bielefeld show higher shares in retail, logistics, and care work.

8 Limitations, Risks & Recommendations

8.1 Limitations

Limitation	Description	Implication
Taxonomy scope	The model is specific to German job titles and KLDB-2010.	Not directly suitable for ISCO, SOC, or other national taxonomies without adaptation.
Taxonomy coverage	Only ~40% of BKZ codes have annotated examples.	The model must generalize to unseen codes via taxonomy names and synonyms.
Class imbalance	The dataset has an extreme long tail.	Head classes may still dominate despite rare-code oversampling.
Sample bias	Green Jobs and AI/KI subsets were selected by external classifiers.	These subsets may introduce systematic selection effects.
Temporal bias	Titles span 2022–2026.	Emerging occupations or shifting terminology after this window may not be covered.
Platform bias	Titles are drawn from online job boards.	Some sectors, such as agriculture, informal employment, and parts of the public sector, may be underrepresented.
Annotation subjectivity	Some titles are genuinely ambiguous.	Fixed mappings encode majority-case assumptions that may not hold for every posting.
Jobbörse calibration snapshot	Generic-title mappings were calibrated against the BA Jobbörse in March 2026.	These distributions can shift over time.
Company names in titles	Some titles may contain employer names.	Employer-specific patterns could leak into the training signal.

Limitation	Description	Implication
Score threshold	The 0.4 cutoff was calibrated on a 100k production sample.	It filters mostly garbage but may occasionally reject unusual legitimate titles.

8.2 Out-of-Scope Use

The dataset and model should not be used to make employment decisions about individuals.

The model is designed for occupational classification of job vacancy titles. It is not a general-purpose labor-market inference model and should not be used outside the German KLDB-2010 context without adaptation and re-evaluation.

8.3 Recommendations

- Monitor the confidence-score distribution in production.
- Spot-check predictions with scores between 0.4 and 0.65.
- Revisit generic-title mappings periodically, because the BA Jobbörse distribution used for calibration may shift over time.
- For internship and working-student titles, evaluate and communicate results at the KLDB-3 level where appropriate.
- Treat Green Jobs and AI/KI performance as domain-specific because the source samples were selected by separate classifiers.