

**Sistem Pembelajaran Adaptif Berbasis Human-in-the-Loop untuk
Ekstraksi Data PDF Template**

TESIS

**Disusun Oleh :
Moh Syaiful Rahman
NPM: 247064518006**



**PROGRAM STUDI MAGISTER TEKNOLOGI INFORMASI
FAKULTAS TEKNOLOGI KOMUNIKASI DAN
INFORMATIKA
UNIVERSITAS NASIONAL
TAHUN 2026**

HALAMAN PENGESAHAN

Sistem Pembelajaran Adaptif Berbasis Human-in-the-Loop untuk Ekstraksi Data PDF Template

Untuk memenuhi sebagai persyaratan memperoleh Gelar Magister Komputer

Disusun oleh:

Moh. Syaiful Rahman

247064518006

Tesis ini telah diuji dan dinyatakan lulus pada

21 Februari 2026

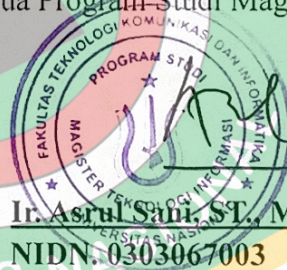
Telah diperiksa dan disetujui oleh:

Dosen Pembimbing,

Mengetahui,
Ketua Program Studi Magister Teknologi Informasi


Dr. Ir. Andrianingsih, S. Kom., MMSI
NIDN. 0303097902


Ir. Asrul Sam, ST., M.T., M.Kom., Ph.D.
NIDN. 0303067003



HALAMAN PENGESAHAN
TESIS

Sistem Pembelajaran Adaptif Berbasis Human-in-the-Loop untuk
Ekstraksi Data PDF Template



Moh Syaiful Rahman

247064518006

Dosen Pembimbing

A handwritten signature in black ink, consisting of a large loop and a horizontal line.

Dr. Ir. Andrianingsih, S. Kom., MMSI

LEMBAR PERSETUJUAN TESIS

Tugas Akhir dengan judul :

Sistem Pembelajaran Adaptif Berbasis Human-in-the-Loop untuk Ekstraksi Data PDF Template

Dibuat untuk melengkapi salah satu persyaratan menjadi Magister Komputer pada Program Studi Magister Teknologi Informasi, Fakultas Teknologi Komunikasi dan Informatika Universitas Nasional. Tesis ini diujikan pada Sidang Tesis Semester Ganjil 2025-2026 pada tanggal 21 Februari Tahun 2026

Dosen Pembimbing

Dr. Ir. Andrianingsih, S. Kom., MMSI
NIDN. 0303097902

Ketua Program Studi

Dr. Asrul Sami, ST., M.T., M.Kom., Ph.D.
NIDN. 0303067003

PERNYATAAN KEASLIAN TESIS

Saya menyatakan dengan sesungguhnya bahwa Tesis dengan judul :

Sistem Pembelajaran Adaptif Berbasis Human-in-the-Loop untuk Ekstraksi Data PDF Template

Yang dibuat untuk melengkapi salah satu persyaratan menjadi Magister Komputer pada Program Studi Magister Teknologi Informasi Fakultas Teknologi Komunikasi dan Informatika Universitas Nasional, sebagaimana yang saya ketahui adalah bukan merupakan tiruan atau publikasi dari Tesis yang pernah diajukan atau dipakai untuk mendapatkan gelar di lingkungan Universitas Nasional maupun perguruan tinggi atau instansi lainnya, kecuali pada bagian – bagian tertentu yang menjadi sumber informasi atau acuan yang dicantumkan sebagaimana mestinya.

Jakarta, 28 Februari 2026



Moh. Syaiful Rahman
NPM: 247064518006

LEMBAR PERSETUJUAN JUDUL YANG TIDAK ATAU YANG DIREVISI

Nama : Moh. Syaiful Rahman

NPM : 247064518006

Fakultas/Akademi : Fakultas Teknologi Komunikasi dan Informatika

Program Studi : Magister Teknologi Informasi

Tanggal Sidang : 21 Februari 2026

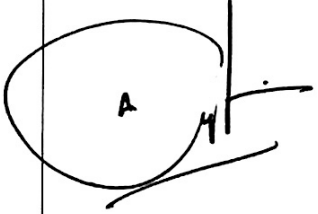
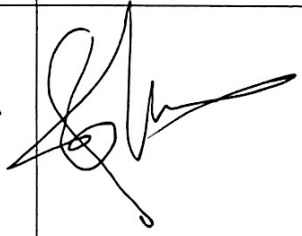
JUDUL DALAM BAHASA INDONESIA :

Sistem Pembelajaran Adaptif Berbasis Human-in-the-Loop
untuk Ekstraksi Data PDF Template

JUDUL DALAM BAHASA INGGRIS :

Adaptive Learning System Based on Human-in-the-Loop
for PDF Template Data Extraction

TANDA TANGAN DAN TANGGAL

Pembimbing 1	Ka. Prodi	Mahasiswa
TGL: 28 Februari	TGL: 28 Februari	TGL: 28 Februari
		

KATA PENGANTAR

Puji syukur penulis panjatkan ke hadirat Allah SWT atas segala rahmat, taufik, dan hidayah-Nya sehingga penulis dapat menyelesaikan tesis yang berjudul “Sistem Pembelajaran Adaptif Berbasis Human-in-the-Loop untuk Ekstraksi Data PDF Template” ini. Tesis ini disusun sebagai salah satu syarat untuk memperoleh gelar Magister Teknologi Informasi pada Program Studi Magister Teknologi Informasi, Fakultas Teknologi Komunikasi dan Informatika, Universitas Nasional.

Penulis menyadari sepenuhnya bahwa penyelesaian tesis ini tidak terlepas dari bantuan, bimbingan, dan dukungan berbagai pihak. Untuk itu, dengan segala kerendahan hati, penulis menyampaikan penghargaan dan ucapan terima kasih yang sebesar-besarnya kepada:

1. Kedua orang tua penulis, Masrikun dan Samaniyah, kakak Jupriadi, adik Ina Fitriani, serta istri tercinta Widi Setyowati, dan buah hati tersayang Prisya Silvana dan Yovika Sklodowska, atas doa yang tidak pernah putus, dukungan moral, dan pengorbanan yang tidak ternilai selama penulis menempuh pendidikan.
2. Jajaran pimpinan Universitas Nasional, mulai dari Rektor, Direktur Program Pascasarjana, hingga Ketua Program Studi Magister Teknologi Informasi, atas segala fasilitas, kebijakan akademik, dan dukungan yang telah diberikan selama masa studi.
3. Dr. Ir. Andrianingsih, S. Kom., MMSI, selaku Dosen Pembimbing Utama, yang dengan penuh kesabaran telah meluangkan waktu untuk memberikan bimbingan, arahan, dan masukan yang sangat berarti dalam setiap tahap penyelesaian tesis ini.
4. Ir. Asrul Sani, ST., M.T., M.Kom., Ph.D. dan Prof. Agung Triayudi, S.Kom, M.Kom, Ph.D., selaku Dosen Penguji, atas kritik dan saran yang tajam dan membangun dalam penyempurnaan tesis ini.
5. Seluruh dosen pengajar Program Studi Magister Teknologi Informasi Universitas Nasional, atas ilmu pengetahuan dan pengalaman akademik yang telah diberikan selama masa perkuliahan.

6. Seluruh staf dan karyawan Fakultas Teknologi Komunikasi dan Informatika Universitas Nasional, atas bantuan dalam kelancaran proses administrasi selama masa studi.
7. Rekan-rekan mahasiswa Program Magister Teknologi Informasi angkatan 2024, atas semangat, diskusi, dan kebersamaan yang menjadi bagian tak terlupakan dari perjalanan studi ini.
8. Semua pihak yang tidak dapat penulis sebutkan satu per satu, yang telah memberikan kontribusi nyata dalam penyelesaian tesis ini.

Penulis menyadari bahwa tesis ini masih memiliki keterbatasan dan kekurangan. Oleh karena itu, kritik dan saran yang konstruktif dari berbagai pihak sangat penulis harapkan demi penyempurnaan ke depan. Penulis berharap tesis ini dapat memberikan kontribusi yang berarti, khususnya dalam pengembangan sistem ekstraksi data berbasis kecerdasan buatan dan penerapan pendekatan Human-in-the-Loop di Indonesia.

Jakarta, 28 Februari 2026

Moh. Syaiful Rahman

247064518006

ABSTRAK

Pada dokumen PDF berbasis template, pemetaan fields sering bergantung pada isyarat tata letak. Perubahan kecil seperti jarak antar-elemen, pergeseran posisi label, atau pergeseran baris pada tabel dapat mengubah hasil pemetaan tersebut. Oleh sebab itu, keluaran ekstraksi dapat berbeda ketika dokumen berasal dari sumber yang berbeda, walaupun jenis template-nya sama. Model pre-trained berukuran besar dapat memberikan akurasi tinggi, tetapi biasanya menuntut komputasi yang besar dan korpus berlabel yang ekstensif, sehingga penerapannya tidak selalu realistis pada lingkungan dengan sumber daya terbatas. Di sisi lain, pemrosesan berbasis aturan relatif ringan dan mudah diaudit, namun sensitif terhadap perubahan tata letak sehingga kinerjanya dapat menurun. Penelitian ini merancang alur pembelajaran adaptif berbasis pendekatan hibrid yang menggabungkan komponen berbasis aturan dengan Conditional Random Fields (CRF) untuk meningkatkan toleransi terhadap variasi tata letak pada kondisi data dan infrastruktur yang terbatas.

Sistem menjalankan beberapa strategi ekstraksi secara paralel dan memilih keluaran berdasarkan kriteria confidence, kemudian mengintegrasikan Human-in-the-Loop (HITL) untuk mendukung perbaikan inkremental. Koreksi pengguna dikonversi menjadi pembaruan pola pada komponen berbasis aturan, sedangkan model CRF dilatih ulang secara berkala dalam batch kecil untuk menangkap fields yang bergantung pada konteks dan sulit ditentukan sepenuhnya melalui aturan. Pendekatan ini tidak dimaksudkan untuk menggantikan salah satu metode, melainkan memanfaatkan aturan untuk interpretabilitas dan pembaruan cepat, serta CRF untuk generalisasi ketika konteks lokal berperan, dengan tetap menjaga kebutuhan deployment yang ringan.

Evaluasi tidak hanya berfokus pada akurasi ekstraksi, tetapi juga pada besarnya usaha koreksi dari pengguna serta efisiensi operasional. Pertimbangan ini penting karena rancangan sistem ditujukan untuk skenario dengan kapasitas anotasi yang terbatas dan ketersediaan komputasi yang tidak selalu memadai. Dievaluasi pada 140 dokumen PDF sintetis di 4 jenis template (Form, Table, Letter, Mixed),

sistem mencapai akurasi 98,61% dengan data pelatihan 25–35 dokumen dan tingkat koreksi pengguna 7,0% (196 dari 2.800 fields), menunjukkan efisiensi pembelajaran 7,55 koreksi per peningkatan poin persentase. Akurasi dihitung pada keluaran ekstraksi sebelum koreksi manual terhadap ground truth, sedangkan koreksi pengguna merepresentasikan akumulasi koreksi sepanjang proses pembelajaran adaptif. Peningkatan performa dari baseline (72,65%) ke sistem adaptif (98,61%) signifikan secara statistik (paired t-test: $t=17,89$, $p<0,001$, Cohen's $d=8,94$). Sistem beroperasi pada perangkat keras CPU-only dengan footprint memori 37–162 MB, ukuran model 1,75 MB, dan waktu pemrosesan median 45 detik per dokumen (termasuk review HITL), sementara waktu ekstraksi inti per dokumen (tanpa HITL) berdasarkan pencatatan waktu sistem berada pada rentang 0,131–0,685 detik (median 0,265 detik)

Secara keseluruhan, hasil ini mengindikasikan bahwa alur kerja hibrid berbasis HITL dapat menyeimbangkan akurasi dengan kebutuhan data dan sumber daya pada lingkungan terbatas. Namun, karena evaluasi menggunakan dokumen sintetis dan koreksi yang disimulasikan, diperlukan validasi lanjutan pada kumpulan dokumen real-world dan studi pengguna sebelum menarik kesimpulan terkait kinerja pada deployment produksi.

Kata Kunci: Conditional Random Fields, Pemrosesan Dokumen, Umpan Balik Pengguna, Skor Kepercayaan, Pelatihan Inkremental

ABSTRACT

Extracting structured fields from template-based PDF documents remains difficult in practice because layouts may vary slightly across sources, while the content is only semi-structured. Although large pre-trained models can deliver strong accuracy, they typically demand substantial compute and sizable labeled corpora, which are not always feasible in resource-constrained settings. At the other end of the spectrum, rule-based parsing is lightweight and transparent, but it can degrade quickly when a document layout shifts. This study therefore explores a hybrid alternative: an adaptive learning workflow that combines rule-based extraction with Conditional Random Fields (CRF) to better accommodate layout variation under limited data and infrastructure

The system runs multiple extraction strategies in parallel and selects outputs using confidence-based criteria, then integrates Human-in-the-Loop (HITL) review to enable incremental improvement. User corrections are converted into pattern updates for the rule-based components, while the CRF models are retrained periodically in small batches to capture context-dependent fields that are difficult to encode as rules. The goal is not to replace one approach with the other, but to use rules for interpretability and fast updates, and CRF for generalization when local context matters, while keeping the deployment footprint lightweight.

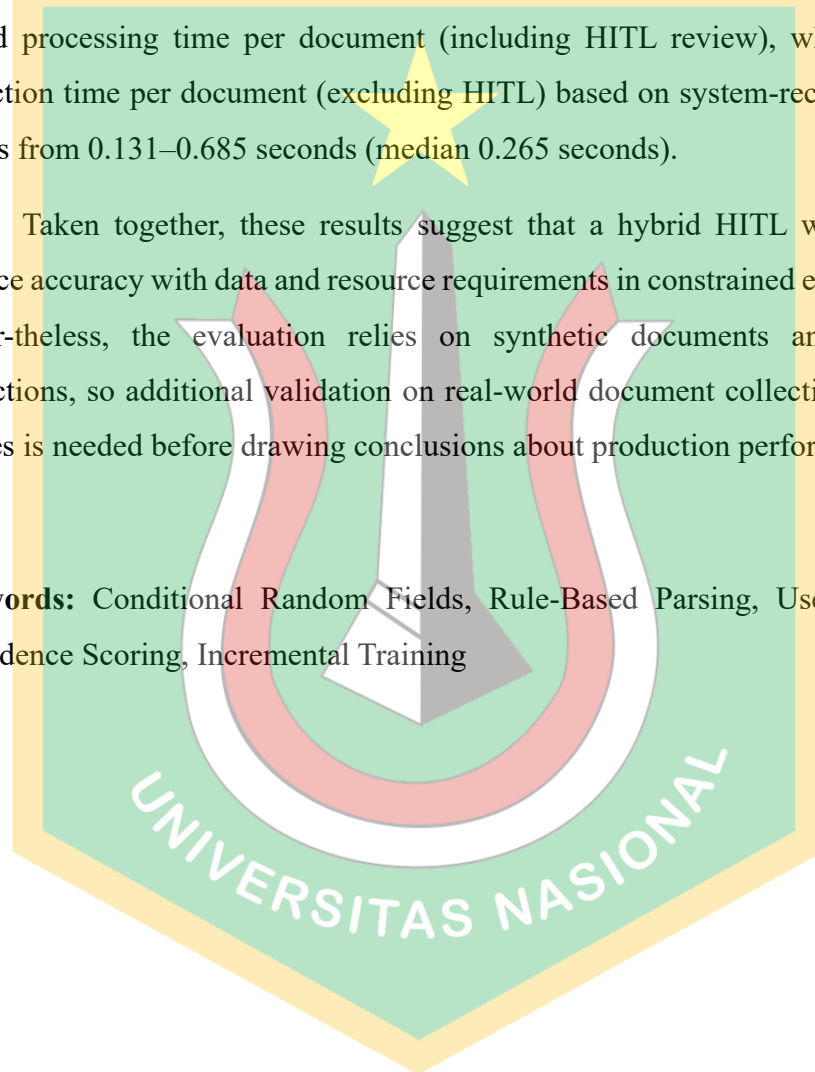
Beyond extraction accuracy, the evaluation also considers user effort and operational efficiency, because the approach targets scenarios with limited annotation capacity and constrained compute resources. The correction-driven update mechanism is intended to reduce labeling overhead by focusing user review on low-confidence fields while preserving a transparent workflow. To support reproducibility of the evaluation, templates and fields are defined consistently across document types, and the reported metrics are computed against ground truth before any user correction is applied.

Evaluated on 140 synthetic PDF documents across 4 template types (Form, Table, Letter, Mixed), the system achieves 98.61% accuracy using 25–35 training documents and a 7.0% user correction rate (196 out of 2,800 fields), corresponding

to a learning efficiency of 7.55 corrections per percentage point improvement. Accuracy is measured on pre-correction extraction output against ground truth, while user corrections represent the accumulated corrections applied throughout the adaptive learning process. The performance improvement from baseline (72.65%) to the adaptive system (98.61%) is statistically significant (paired t-test: $t=17.89$, $p<0.001$, Cohen's $d=8.94$). The system operates on CPU-only hardware with 37–162 MB memory footprint, 1.75 MB total model size, and median 45 seconds end-to-end processing time per document (including HITL review), while the core extraction time per document (excluding HITL) based on system-recorded timing ranges from 0.131–0.685 seconds (median 0.265 seconds).

Taken together, these results suggest that a hybrid HITL workflow can balance accuracy with data and resource requirements in constrained environments. Nevertheless, the evaluation relies on synthetic documents and simulated corrections, so additional validation on real-world document collections and user studies is needed before drawing conclusions about production performance.

Keywords: Conditional Random Fields, Rule-Based Parsing, User Feedback, Confidence Scoring, Incremental Training



DAFTAR ISI

BAB 1	PENDAHULUAN	1
1.1	Latar Belakang	1
1.2	Rumusan Masalah	2
1.3	Tujuan	3
1.4	Batasan Penelitian	3
1.5	Manfaat Penelitian	4
1.5.1	Manfaat Teoritis	4
1.5.2	Manfaat Praktis	4
1.6	Metodologi Penelitian	5
1.7	Sistematika Penulisan	5
BAB 2	TINJAUAN PUSTAKA	7
2.1	Dokumen PDF Semi-Terstruktur	7
2.1.1	Pengertian dan Karakteristik Dokumen PDF	7
2.1.2	Klasifikasi Dokumen Berdasarkan Struktur	7
2.1.3	Karakteristik Form PDF Semi-Terstruktur	8
2.2	Teknik Ekstraksi Data dari Dokumen PDF	9
2.2.1	Pendekatan Berbasis Aturan	9
2.2.2	Pendekatan Berbasis Machine Learning	9
2.2.3	Model Pre-trained untuk Pemahaman Dokumen	10
2.3	Sistem Human-in-the-Loop	11
2.3.1	Paradigma Human-in-the-Loop	11
2.3.2	Strategi Pembelajaran Efisien	11
2.3.3	Kolaborasi Manusia-AI	12
2.3.4	Polarisasi Pendekatan	12

2.3.5	Temuan Penelitian HITL	12
2.4	Pembelajaran Adaptif dan Continual Learning.....	13
2.4.1	Continual Learning.....	13
2.4.2	Karakteristik Continual Learning.....	13
2.4.3	Aplikasi dalam Ekstraksi Dokumen.....	13
2.5	Adaptive Strategy Selection	13
2.5.1	Hybrid Extraction Systems	13
2.5.2	Adaptive Weighting Mechanism.....	14
2.6	Analisis Komparatif.....	14
2.7	Kesenjangan Penelitian	17
2.7.1	Kesenjangan Sumber Daya	17
2.7.2	Peran HITL Terbatas dalam Model Besar	18
2.7.3	Keterbatasan Sistem Transparan	18
2.7.4	Arsitektur Hibrid Adaptif yang Efisien	18
2.8	Positioning Penelitian Ini.....	19
BAB 3	METODOLOGI PENELITIAN	21
3.1	Desain Penelitian	21
3.1.1	Pendekatan Penelitian	21
3.1.2	Kerangka Design Science Research.....	21
3.1.3	Tahapan Penelitian	23
3.1.4	Pemilihan Metode	24
3.2	Arsitektur Sistem	24
3.2.1	Gambaran Umum Arsitektur	25
3.2.2	Alur Data Sistem.....	26
3.2.3	Model Data Sistem.....	31
3.2.4	Komponen Implementasi Sistem	33
3.3	Strategi Ekstraksi.....	37
3.3.1	Strategi Berbasis Aturan.....	37
3.3.2	Strategi Berbasis CRF	38
3.3.3	Mekanisme Pemilihan Strategi	39

3.4	Mekanisme Human-in-the-Loop.....	44
3.4.1	Komponen HITL.....	44
3.4.2	Integrasi Umpan Balik	45
3.4.3	Mekanisme Pembelajaran Adaptif	45
3.4.4	Proses Hibrid Ekstraksi dan Pembelajaran	46
3.5	Metrik Evaluasi	47
3.5.1	Metrik Kinerja Standar.....	47
3.5.2	Metrik Pembelajaran Adaptif.....	48
3.5.3	Metrik Efisiensi Sumber Daya	49
3.6	Prosedur Eksperimen	49
3.6.1	Dataset.....	49
3.6.2	Prosedur Pengujian.....	51
3.6.3	Detail Implementasi	55
	Pertimbangan Etis	56
3.7	Validitas Penelitian.....	56
3.7.1	Validitas Internal	56
3.7.2	Validitas Konstruk.....	57
3.7.3	Validitas Eksternal.....	58
BAB 4	HASIL.....	61
4.1	Implementasi Sistem	61
4.1.1	Antarmuka Pengguna	61
4.1.2	Workflow Pengguna.....	81
4.2	Hasil Eksperimen	85
4.2.1	Kinerja Keseluruhan.....	85
4.2.2	Contoh Ekstraksi Representatif.....	87
4.2.3	Analisis Metrik Detail	94
4.2.4	Ablation Study	98
4.2.5	Analisis Reduksi Error	100
4.2.6	Analisis Pembelajaran Adaptif.....	102
4.2.7	Ringkasan Hasil Eksperimen	104

BAB 5 PEMBAHASAN.....	106
5.1 Efektivitas Pendekatan Hibrid.....	106
5.1.1 Kekuatan Komplementer	106
5.1.2 Validasi Hipotesis.....	107
5.2 Efektivitas Pembelajaran HITL	107
5.2.1 Minimal User Burden.....	107
5.2.2 High Learning Efficiency.....	108
5.2.3 Fast Convergence	108
5.2.4 Incremental Learning	109
5.3 Kemampuan Generalisasi	109
5.3.1 Consistent Performance	109
5.3.2 Template Adaptability	110
5.3.3 Scalability.....	110
5.3.4 Error Analysis.....	111
5.4 Implikasi Praktis	111
5.4.1 Potensi Production Readliness	111
5.4.2 Business Value	112
5.4.3 Deployment Scenarios	113
5.5 Perbandingan dengan State-of-the-Art.....	113
5.5.1 Trade-off Utama	114
5.5.2 Alignment dengan Penelitian Terkini.....	114
5.6 Keterbatasan Penelitian.....	115
5.6.1 Keterbatasan Saat Ini.....	115
5.6.2 Arah Penelitian Masa Depan.....	115
5.7 Positioning dalam Lanskap Penelitian	116
5.7.1 Complementary Solution	117
5.7.2 Target Scenarios	117
5.7.3 Kontribusi Unik.....	118
BAB 6 KESIMPULAN DAN SARAN	119
6.1 Kesimpulan	119

6.1.1	Pencapaian Utama.....	119
6.1.2	Menjawab Kesenjangan Penelitian	120
6.1.3	Kontribusi Penelitian.....	120
6.1.4	Implikasi Penelitian.....	122
6.1.5	Dampak Lebih Luas	123

6.2 Keterbatasan Penelitian..... 123

6.2.1	Validitas Eksternal.....	123
6.2.2	Keterbatasan Metodologi	124
6.2.3	Implikasi Keterbatasan.....	124

6.3 Saran..... 125

6.3.1	Pengembangan Sistem	125
6.3.2	Penelitian Lanjutan.....	126
6.3.3	Implementasi Praktis.....	127
6.3.4	Future Vision.....	128
6.3.5	Penutup.....	128



DAFTAR GAMBAR

Gambar 3-1 System architecture hybrid extraction approach.	25
Gambar 3-2 Data Flow Diagram Level 0 (Context Diagram).....	26
Gambar 3-3 Data Flow Diagram Level 1	27
Gambar 3-4 UML Use Case Diagram sistem ekstraksi data	29
Gambar 3-5 UML Activity Diagram alur kerja sistem dari input dokumen hingga pembelajaran adaptif.	29
Gambar 3-6 Sequence Diagram Adaptive Learning.....	30
Gambar 3-7 Entity Relationship Diagram	31
Gambar 3-8 Strategy selection flowchart with confidence-based comparison and threshold checking.....	40
Gambar 3-9 HITL feedback flow with incremental learning from user corrections.....	45
Gambar 4-1 Implementasi Template Management - Upload Form dan List Template	62
Gambar 4-2 Implementasi Template Management - Preview dan Field Configuration	63
Gambar 4-3 Implementasi Document Extraction - Upload dan Results List	64
Gambar 4-4 Implementasi Document Extraction - Detail Results dan Feedback Form.....	65
Gambar 4-5 Implementasi Model Training - Manual dan Incremental Training Interface	66
Gambar 4-6 Implementasi Dashboard - Overview Tab dengan Performance Metrics.....	68
Gambar 4-7 Implementasi Dashboard - Adaptive Learning Progress Chart dan Key Insights.....	70
Gambar 4-8 Implementasi Dashboard - Field Analysis: Field-Level Accuracy dan Confidence Indicators.....	71
Gambar 4-9 Implementasi Dashboard - Field Analysis: Detailed Evaluation Metrics dengan Precision, Recall, F1-Score	72

Gambar 4-10 Implementasi Dashboard - Strategies: Strategy Distribution, Performance, dan Confidence Trends	73
Gambar 4-11 Implementasi Dashboard - Strategies: Ablation Study untuk Strategy Comparison	74
Gambar 4-12 Implementasi Dashboard - Performance: Extraction Time Statistics dan Trends.....	75
Gambar 4-13 Implementasi Dashboard - Error Patterns: Analysis untuk Identifikasi Problematic Fields.....	77
Gambar 4-14 Implementasi Dashboard - Research Metrics: HITL Metrics dan Adaptive Learning Status	78
Gambar 4-15 Implementasi Dashboard - Research Metrics: Incremental Learning Progress dan Baseline Comparison.....	80
Gambar 4-16 Error reduction comparison.....	101
Gambar 4-17 Learning curves showing accuracy improvement across batches for all template types.....	102
Gambar 4-18 Learning efficiency showing relationship between corrections and accuracy improvement.....	103



DAFTAR TABEL

Tabel 2-1 Perbandingan Penelitian Terdahulu.....	15
Tabel 4-1 Perbandingan Kinerja Baseline vs. Adaptive per Template	85
Tabel 4-2 Contoh Ekstraksi Form Template - Perfect Case	87
Tabel 4-3 Contoh Ekstraksi Form Template - Corrected Case.....	88
Tabel 4-4 Contoh Ekstraksi Table Template - Perfect Case	89
Tabel 4-5 Contoh Ekstraksi Table Template - Corrected Case.....	89
Tabel 4-6 Contoh Ekstraksi Letter Template - Perfect Case	90
Tabel 4-7 Contoh Ekstraksi Letter Template - Corrected Case.....	91
Tabel 4-8 Contoh Ekstraksi Mixed Template - Perfect Case	92
Tabel 4-9 Contoh Ekstraksi Mixed Template - Corrected Case.....	93
Tabel 4-10 Metrik Evaluasi Detail (Baseline vs. Adaptive).....	94
Tabel 4-11 Statistical Significance Tests: Baseline vs Adaptive System ..	96
Tabel 4-12 5-Fold Cross-Validation Results: Adaptive System.....	97
Tabel 4-13 Ablation Study: Component Contribution Analysis	98
Tabel 4-14 Tingkat Koreksi Pengguna per Template	104
Tabel 4-15 Ringkasan Hasil Eksperimen Utama	104
Tabel 5-1 Perbandingan Efisiensi Sumber Daya Sistem.....	111

UNIVERSITAS NASIONAL