

以下の論文は、ChatGPT をカスタマイズした「技術士二次試験三上塾 AI 論文作成ツール」

<https://suninfo.co.jp/g2gpt/>

が作成した「概要論文」なので、一般的な合格論文と異なります。

勉強の進め方としては、以下の資料を元に論文の方向性を掴みつつ、概要表や出典などを参考に以下の概要論文を修正すると短時間で効率的な勉強ができます。以下の出典の URL をクリックしてもサイトが表示されない場合は、出典名をネット検索してください。

以下の論文では、観点と課題を章の最初に書いていますが、課題などを最後に書きたい場合は、テンプレート（プロンプト）を変更すれば最後に書くことが可能です。

設問は以下のサイトをご覧ください。

https://www.engineer.or.jp/c_categories/index02022.html

=====

以下では、**選択科目Ⅲで問われる「問題解決能力・課題遂行能力」**、特に「複合的な問題を把握→多面的な観点から技術課題を抽出→解決策→実行後のリスク→対策」という論理連鎖を重視して整理します。選択科目Ⅲでは、単なる専門知識の列挙ではなく、この一連の流れを**論理的かつ合理的に説明できる指導的技術者であることを答案で示すことが重要**です。

1. 設問テーマの社会目的

この設問の社会目的は、

データとアルゴリズムを活用して、利用者と適切な相手・リソースを効率的かつ公正・安全に結び付け、社会に存在する人的・物的リソースの有効活用と利用者便益の向上を実現すること

です。

重要なのは「高精度な AI を作ること」自体を目的にしないことです。

マッチングサービスでは、デジタル技術によって市場へのアクセスを拡大できる一方、プラットフォームの透明性・公正性も社会的な論点となっています。（[経済産業省](#)）

したがって、

社会目的 → マッチング成立の質・量を向上 → そのための技術課題

という順番で考えます。

2. 出題者の意図

この問題は「レコメンド技術を知っていますか」という知識問題ではありません。

令和8年度の選択科目Ⅲは、社会的ニーズや技術進歩から生じる複合的問題について、選択科目に関わる観点から技術課題を抽出し、多様な視点から分析して問題解決手法と遂行方策を提示できるかを問う試験です。

本問では特に、次の能力を確認していると考えます。

設 問	主に確認される能力
(1)	具体的サービスをモデル化し、データ・アルゴリズム・オペレーション等から複合的問題を分析できるか
(1)	「問題」と「技術課題」を区別し、多面的な観点から3課題を抽出できるか
(2)	重要課題を評価・選定し、専門知識を使って複数の解決策を構築できるか
(3)	解決後にも残る・新たに発生するリスクを予測し、さらに対策できるか

特に「データ、アルゴリズム、オペレーション等に着目」と書かれている点がヒントです。

つまり、出題文自体が多面的分析の切り口をかなり明確に与えています。

3. 重要な箇所と注意すべき箇所

最重要①「具体的なマッチング系サービスを想定」

抽象的な「マッチングシステム」のまま論じると答案が薄くなります。

今回は、IT人材と企業案件を結ぶ人材マッチングサービスを想定します。

これなら、

スキルデータ → 検索・推薦 → 双方向マッチング → 面談・契約 → 評価フィードバック

という設問①～⑤との対応が明確です。

最重要②「多面的な観点」

単に課題を3個書くのではなく、

観点 → 問題分析 → 技術課題

の順にします。

マニュアルでも、骨子を「現状・問題・求められる技術課題→観点・技術課題の抽出→解決策→解決後のリスク→対応策」という流れで整理しています。

今回は設問文を利用して、

データ品質／アルゴリズム／オペレーション

の3観点にすると整理しやすいです。

最重要③「技術課題」

ここは「〇〇が不足している」という**問題表現**ではなく、

〇〇の高度化

〇〇の確保

〇〇の構築

という**達成すべき状態**で表現します。

例えば、

×「登録スキル情報が不正確であることが課題」

より、

○「信頼性の高い人材・案件データの確保が課題」

とします。

最重要④「すべての解決策を実行しても」

ここは非常に重要です。

解決策1だけの副作用を書くのではなく、

解決策1+2+3を全部実施した世界を想像する → それでも何が起こるか

と考えます。

マニュアルでも、骨子表は「すべての解決策を実行した上で生じる波及効果・懸念事項」を考える構造になっています。

4. 論文作成上のポイント

この問題で最も作りやすい論理軸は、

「マッチング精度」だけを追わず、「信頼できるマッチング」を実現する

ことです。

具体的には、

正しいデータ → 適切なアルゴリズム → 安全な運用 → 信頼されるサービス

という一本のストーリーにします。

さらに情報工学部門らしさを出すため、

Embedding、Learning to Rank、Two-Tower、A/B テスト、MLOps、Concept Drift、Explainable AI、RBAC、多要素認証、異常検知

などを「名前だけ」でなく、何のために使うかまで書きます。

マニュアルでも、専門性を示す重要キーワードを論文へ埋め込み、現状・課題・解決策を具体化することが重視されています。

5. 不合格者と合格者の思考の違い

項目	不合格になりやすい思考	合格答案を作る思考
テーマ理解	AI・マッチングの知識を書く	社会目的から逆算する
サービス	抽象的なマッチングサービス	IT 人材マッチング等を具体化
問題	課題と問題を混同	問題→原因→技術課題と考える
観点	思いついた 3 課題	データ・アルゴリズム・運用等で分解
課題表現	データ不足が課題	データ品質の確保が課題
最重要課題	何となく選択	社会目的への影響度で評価
解決策	AI を導入する	Embedding、LTR、MLOps 等まで具体化

項目	不合格になりやすい思考	合格答案を作る思考
専門用語	用語を大量に列挙	用語＋目的＋実装方法を書く
リスク	解決前から存在する問題を書く	全解決策実施後の新たなリスクを考える
答案全体	各章がバラバラ	現状→課題→解決策→リスクを一本化

6. 概要表

今回は、IT エンジニア・IT 企業案件の人材マッチングサービスとして骨子を構築します。

項目	内容
設問の概要（現状・問題）	デジタル技術によりマッチングサービスが普及している。人材マッチングでは大量の人材・案件から適切な組合せを短時間で生成できる一方、自己申告データの不正確さ、推薦アルゴリズムの偏り、虚偽登録・不正利用等により、ミスマッチや利用者の信頼低下が発生し得る。
社会目的	IT 人材と企業案件を適切かつ安全に結び付け、人的資源を有効活用するとともに、双方の満足度と継続利用率を高める。
想定サービス	IT エンジニアと企業案件を結ぶ人材マッチングサービス。
機能概要	人材側はスキル、経験、希望職種、勤務地等を登録し、企業側は必要スキル、業務内容、報酬等を登録する。検索・推薦・双方向マッチングにより候補を提示し、面談・契約を経て、終了後の評価を推薦モデルへフィードバックする。
技術課題①	観点：データ品質。課題：信頼性の高い人材・案件データの確保。自己申告ではスキル表記の揺れ、誇張、情報の陳腐化が生じ、入力データが不正確なら推薦精度も低下する。
技術課題②	観点：アルゴリズム。課題：公平性・説明可能性を備えたマッチング精度の高度化。成約履歴だけで学習すると人気人材・企業への偏重やコールドスタートが生じ、多様な候補を発見できない。
技術課題③	観点：オペレーション・安全性。課題：信頼性の高い取引環境の構築。なりすまし、虚偽案件、不正アクセス、個人情報漏えい等が発生すると利用者の安全とサービスへの信頼を損なう。
最重要課題	公平性・説明可能性を備えたマッチング精度の高度化。マッチングの中核機能であり、精度・公平性が低ければ良質なデータや安全な運用基盤を整えても成立率・満足度が向上せず、社会目的を達成できないため。
解決策①	意味ベース推薦の高度化。スキル・職務経歴・案件要件を Embedding でベクトル化し、Two-Tower モデル等で候補生成する。Learning to Rank でスキル適合度、希望条件、評価履歴等を再ランキングする。

項目	内容
解決策②	双方向最適化と公平性確保。 人材側・企業側双方の選好を予測し、成立確率を目的関数としてランキングする。Exposure 等の公平性指標を監視し、人気偏重を抑制する。推薦理由を提示し説明可能性も確保する。
解決策③	継続学習基盤の構築。 応募、面談、成約、評価をフィードバックし、A/B テストで推薦モデルを評価する。MLOps により精度、公平性、Concept Drift を監視し、必要に応じ再学習する。
リスク	推薦精度向上と継続学習を進めるほど、過去のクリック・成約実績が次の推薦へ反映され、人気人材・企業がさらに露出する フィードバックループ が形成される。結果として特定属性・職種が過小推薦されるアルゴリズムバイアスが増幅するリスクがある。
対策	Accuracy や成約率だけでなく、Exposure、Demographic Parity 等の公平性指標を継続監視する。リランキングで露出制約を設定し、探索枠を確保する。また、モデル・データのバージョン管理、定期的なバイアス監査、Human-in-the-Loop による異常推薦のレビューを実施する。

AI リスクは技術だけでなく組織的に継続管理する必要があります。NIST AI RMF も AI システムの設計・開発・利用・評価を通じたリスク管理を目的としており、日本の AI 事業者ガイドラインも 2026 年に第 1.2 版へ更新されています。[\(NIST\)](#)

また、人材情報は個人情報も多く扱うため、個人情報保護とセキュリティも背景知識として押さえておくべきです。個人情報保護委員会は個人情報の適正な取扱いの具体的指針を示しており、IPA「情報セキュリティ 10 大脅威 2026」でも不正ログインや個人情報窃取等が個人向け脅威として挙げられています。[\(警察庁\)](#)

参考出典

1. 経済産業省「AI 事業者ガイドライン（第 1.2 版）」
[AI 事業者ガイドライン 第 1.2 版](#)
2. 個人情報保護委員会「個人情報の保護に関する法律についてのガイドライン（通則編）」
[個人情報保護法ガイドライン（通則編）](#)
3. 経済産業省「デジタルプラットフォーム」
[経済産業省 デジタルプラットフォーム政策](#)
4. 独立行政法人情報処理推進機構「情報セキュリティ 10 大脅威 2026」
[情報セキュリティ 10 大脅威 2026](#)
5. National Institute of Standards and Technology, “Artificial Intelligence Risk Management Framework (AI RMF 1.0)”
[NIST AI Risk Management Framework 1.0](#)

6. 経済産業省「特定デジタルプラットフォームの透明性及び公正性の向上に関する法律」

[デジタルプラットフォーム取引透明化法](#)

7. 約 1800 文字の論文

以下は、上記骨子を指定された章立てに落とした**答案例**です。マニュアルの「現状・問題→観点・技術課題→解決策→解決後のリスク→対応策」という骨子構造を踏襲しています。

1. 多面的な観点からの技術課題の抽出

0) 想定した具体的なマッチング系サービスと機能の概要

IT エンジニアと企業案件を結ぶ人材マッチングサービスを想定する。人材側はスキル、職務経験、希望職種等を、企業側は必要スキル、業務内容、報酬等を登録する。システムは条件検索やレコメンドにより候補を抽出し、双方の選好を考慮して組合せを提示する。その後、面談・契約を行い、成約結果や双方の評価を推薦モデルにフィードバックする。

1) 信頼性の高い人材・案件データの確保

データ品質の観点で考えれば、信頼性の高い人材・案件データの確保が課題である。具体的には、スキルや職歴等を自己申告だけに依存すると、表記揺れ、誇張、情報の陳腐化が生じる。また、同一技術でも登録者ごとに粒度が異なれば、検索や推薦の入力データの品質が低下し、ミスマッチにつながる。このため、データを標準化し、正確性、完全性、最新性を確保する必要がある。

2) 公平性・説明可能性を備えたマッチング精度の高度化

アルゴリズムの観点で考えれば、公平性・説明可能性を備えたマッチング精度の高度化が課題である。具体的には、過去のクリックや成約実績だけで推薦すると、実績の多い人材や企業に推薦が集中する人気バイアスが生じる。また、新規利用者には履歴がなくコールドスタート問題が発生する。このため、双方の希望を考慮しつつ、精度、公平性、多様性を確保した推薦が必要となる。

3) 信頼性の高い取引環境の構築

オペレーション・安全性の観点で考えれば、信頼性の高い取引環境の構築が課題である。具体的には、なりすまし、虚偽案件、不正アクセス等が発生すれば、個人情報や取引情報が悪用され、利用者の信頼を失う。このため、本人確認、アクセス制御、不正検知等を組み合わせ、取引の安全性を確保する必要がある。

2. 抽出した技術課題のうち最も重要と考える技術課題と解決策

私は「公平性・説明可能性を備えたマッチング精度の高度化」が最も重要な課題と考える。理由は、マッチングがサービスの中核機能であり、適切な組合せを生成できなければ、良質なデータや安全な運用基盤を整備しても、成約率や利用者満足度の向上につながらないからである。

1) 意味ベース推薦の高度化

人材と案件の意味的な適合度を高めることが重要である。具体的には、スキル、職務経歴、案件要件等のテキストを Embedding によりベクトル化し、Two-Tower モデルで候補を高速に抽出する。さらに Learning to Rank を用いて、スキル適合度、希望条件、過去の評価等を特徴量として候補を再ランキングする。

2) 双方向最適化と公平性の確保

人材側と企業側の双方が希望する組合せを生成することが重要である。具体的には、双方の応募・承諾履歴から選好を推定し、成約確率を目的関数としてランキングする。また、特定の人材や企業への過度な集中を防止するため、Exposure 等の公平性指標を監視し、ランキングにより多様性を確保する。併せて推薦理由を提示し、説明可能性を高める。

3) 継続学習基盤の構築

利用実績を用いて継続的に推薦性能を改善することが重要である。具体的には、応募、面談、成約、評価等をフィードバックデータとして収集し、A/B テストによりモデルを評価する。また、MLOps を導入し、精度や公平性、Concept Drift を継続監視して、性能低下時にはモデルを再学習する。

3. 解決策を実行しても新たに生じうるリスクと対策

1) リスク

リスクとしては、推薦結果のフィードバックループによるアルゴリズムバイアスの増幅がある。具体的には、高精度化と継続学習を進めるほど、過去に多く推薦された人材・企業ほどクリックや成約データを獲得し、それが次の学習データとなる。その結果、人気のある候補がさらに優先され、新規利用者や特定属性の候補が過小評価されるおそれがある。

2) 対策

対策としては、公平性を含む AI ガバナンスの構築がある。具体的には、成約率や Accuracy だけでなく、Exposure 等の公平性指標を継続監視する。また、リランキングに露出制約を設定するとともに、新規候補を提示する探索枠を設ける。さらに、モデルと学習データをバージョン管理し、定期的なバイアス監査を実施する。異常な推薦については Human-in-the-Loop で技術者が確認し、必要に応じてロールバックや再学習を行う。これにより、推薦精度と公平性を両立し、利用者が継続して信頼できるマッチングサービスを実現する。

この答案で特に押さえておきたい点

この答案の中心は**「AI を使えばマッチングできる」ではありません**。

採点者に見せたいのは、

問題： データ不備・人気偏重・不正利用がある

↓

課題： データ品質／アルゴリズム／オペレーションから 3 課題を抽出

↓

評価： 社会目的への影響からアルゴリズムを最重要と判断

↓

解決： 候補生成→双方向最適化→継続学習

↓

新リスク： 継続学習によってフィードバックループが発生

↓

対策： 公平性監視・探索・監査・Human-in-the-Loop

という因果関係の通った問題解決プロセスです。

選択科目Ⅲでは、この「現状・問題→技術課題→解決策→リスク→対策」が論理的につながっているかを最後に骨子表へ戻って確認することが非常に重要です。これは、マニュアルが求める「問題解決のための課題とその遂行を論理的かつ合理的に説明する」という方向性にも合致します。

以上