

TECHNICAL SCIENCES

AI-POWERED TASK ASSIGNMENT RECOMMENDATION SYSTEM FOR SOFTWARE DEVELOPMENT TEAMS

Turysbay D.,

Master Student, School of Information Technology and Applied Mathematics, SDU University, Kaskelen, Kazakhstan

<https://orcid.org/0009-0003-6366-4322>

Rysbayeva G.

PhD in Cybersecurity, Senior lecturer, School of Information Technology and Applied Mathematics SDU University, Kaskelen, Kazakhstan

<https://orcid.org/0000-0002-7331-353X>

<https://doi.org/10.5281/zenodo.19893261>

Abstract

Assigning tasks to the right developers in a software project is essential for timely delivery, but this process usually depends on manual judgment and management intuition. In this work, we introduce SmartAssign—a machine learning-driven recommendation platform—through a principled multi-factor scoring method. The proposed hybrid algorithm uses four complementary dimensions: skill adjustment (40%), previous performance mode (30%), current ability distribution (20%), and qualifications suitable for complexity (10%), to generate a list of ranked candidates, accompanied by manual explanations. We have established a complete prototype, including an offline training pipeline, a RESTful reasoning service, and a browser-based management dashboard. Eight widely used classification models are benchmarked on a corpus of 500 synthetic assignment records; the highest-performing random forest classifier achieved an accuracy of 82.7% with 0.891 ROC-AUC, while XGBoost followed closely with an accuracy of 81.3% and 0.885 ROC-AUC. Compared with mainstream project tracking platforms that only consider developer availability, SmartAssign provides context-aware, evidence-based recommendations, considering proficiency depth, cognitive burden, and outcome history. Each recommendation is accompanied by a transparent theoretical basis, empowering team leads to make well-founded allocation choices.

Keywords: task assignment, recommendation systems, machine learning, software project management, explainable AI

Introduction.

Assigning work items to various contributors is one of the most important and error-prone activities in software project management. When a new task enters the backlog, the responsible manager must weigh a series of factors: the candidates' technical track record, their qualifications relative to the difficulty of the task, the amount of work they have undertaken, and their reliability in completing similar tasks in the past. In practice, the vast majority of these decisions are made intuitively, which introduces inefficiency into the system. Industry surveys show that about 66% of software companies have missed their original timetable targets, while close to 45% have exceeded their planned budgets [1].

Contemporary project management platforms—Jira, Asana, Monday.com, Trello, and their peers—provide powerful tracking and workflow orchestration, but do not provide substantive guidance on who should perform a given task. Assignment is regarded as a simple slot-filling exercise: if the developer appears available, you can drag the task to their board. The rich, multi-dimensional relationship between what the task needs and what the developer can provide is largely ignored.

This article addresses the following research question: Can machine learning be used to recommend optimal developer-task pairings in software projects to improve the distribution success rate while maintaining a balanced workload across the entire team?

Several intertwined challenges make this issue extraordinary. First, skill matching is inherently subtle: a task may call for three different abilities at different ability thresholds, and shallow keyword overlap can underestimate or exaggerate true suitability. Second, workload fairness is important—transferring highly complex projects to the best performers accelerates burnout, while junior staff are underutilized. Third, historical signals are informative: past completion times, quality ratings, and redistribution events encode patterns that pure profile matching misses. Fourth, transparency is a prerequisite for adoption: managers are unlikely to follow opaque algorithmic instructions without understanding the reasoning. Fifth, the delay must be low: agile teams make allocation decisions at time-limited ceremonies such as sprint planning, so recommendations must be reached within a few seconds.

The task allocation problem admits several complementary formalizations. Through the optimization lens, it is a multi-objective problem that simultaneously focuses on maximum skill utilization, minimum workload variance, and maximum predicted success. Through the recommendation system lens, it is similar to item recommendation—developers are the "items" recommended to task "users" on the basis of historical interaction signals. Through the supervised learning lens, each developer-task pair can be labeled as a likely success or failure, which makes standard classification

machinery applicable. SmartAssign unifies all three perspectives within one framework.

Specifically, our contributions are fourfold: (1) a purpose-built hybrid scoring algorithm for software task allocation that combines skill proficiency depth, result history, real-time capacity, and level alignment; (2) a systematic empirical comparison of eight classification algorithms applied to the prediction of allocation outcomes; (3) a domain-informed feature engineering model encoding twelve signals derived from software development practice; and (4) an explanatory layer that produces natural language rationales together with SHAP-based feature attributions, packaged in a deployable full-stack prototype.

The rest of this article proceeds as follows. Section 2 surveys prior work on task assignment, recommendation systems, and interpretable AI in software engineering. Section 3 introduces the hybrid algorithm, feature design, model selection, and system architecture in detail. Section 4 reports the experimental results and discusses their implications. Section 5 concludes and outlines directions for future investigation.

Literature Review.

Task Allocation in Software Engineering. Distributing development work has been examined from diverse angles over the past two decades. Early contributions cast the problem as a constraint-satisfaction or linear-programming exercise [2], treating developers as largely interchangeable units whose availability is the binding constraint. While computationally elegant, such formulations strip away the heterogeneity of human expertise. Subsequent skill-oriented methods introduced similarity measures—the Jaccard coefficient, cosine distance—to compare task requirements against developer profiles, yet they remained static: neither current workload nor track-record information entered the matching score. Qualitative field studies confirm that practicing managers weigh a far richer set of considerations, including domain familiarity, mentoring opportunities, and interpersonal compatibility [3], underscoring the gap between existing models and real decision-making.

Recommender Systems and Predictive Models. The broader recommendation system literature provides three basic paradigms [4], [5]: collaborative filtering, which infers preferences from similar user behavior but struggles in cold-start scenarios; content-based filtering, which uses feature similarity between items and user profiles; and hybrid strategies, which combine signals from the two families to mitigate the disadvantages of each. SmartAssign adopts a hybrid philosophy that pairs content-based skill comparison with collaborative history-result patterns and contextual workload status.

In the field of software engineering, automatic bug triage pipelines use text-based methods to route defect reports to likely fixers. Classifiers—usually Naïve Bayes or support vector machines—are trained on bug report tokens [6]–[8]. Developer recommendation engines mine version control logs and code ownership diagrams to suggest reviewers or collaborators [9], [10],

[16]. Our work goes beyond expertise extraction by integrating capacity constraints and multi-metric performance history into the ranking function.

The growing emphasis on interpretable artificial intelligence [11], [12] is particularly relevant here: empirical evidence shows that practitioners are more willing to act on algorithmic recommendations when accompanied by understandable rationales. Techniques such as SHAP and LIME provide post-hoc explanations for complex ensemble outputs. This is the approach we adopt in SmartAssign.

Methods.

Problem Formulation. We frame developer-task assignment as a joint classification and ranking task. Let

$D = \{d_1, d_2, \dots, d_n\}$ denote the pool of developers, each characterized by a skill profile, a seniority grade, and a real-time workload snapshot. Given an incoming task t described by its required skills, complexity tier, estimated effort, and priority level, together with a historical log $H = \{(d_i, t_j, outcome)\}$ of past assignments and their results, the goal is to estimate, for every candidate pair (d_i, t) , the probability that the assignment will succeed, and to rank candidates by that probability. An assignment qualifies as successful when (a) the task finishes within $\pm 20\%$ of the time estimate, (b) the delivered quality score is at least 4.0 out of 5.0, and (c) no reassignment occurs.

The recommendation engine computes a weighted composite of four orthogonal sub-scores. 1) Skill Match Score (Weight: 40%): Developer competencies are encoded as vectors of proficiency ratings (scale 1–5), while task requirements are represented as binary need indicators. The proficiency-weighted alignment score is:

$$S_{\text{skill}} = \frac{1}{|S_{\text{req}}|} \sum_{s \in S_{\text{req}}} \frac{\min(\text{prof}_{\text{dev}}(s), \text{prof}_{\text{req}}(s))}{\text{prof}_{\text{req}}(s)} \quad (1)$$

where S_{req} is the set of skills the task demands.

2) Performance History Score (Weight: 30%): Three retrospective indicators are fused:

$$\text{Sperf} = 0.4 \times \text{success rate} + 0.3 \times \text{quality} + 0.3 \times \text{time accuracy} \quad (2)$$

Developers who lack prior records receive a neutral default of 0.6 to mitigate cold-start bias.

Feature Engineering. Twelve predictive features were constructed across three groups. Core features: skill match score (cosine-based, range 0–1), workload ratio, grade correspondence (categorical), complexity tier (1–5), and priority encoding (1–3). Historical features: cumulative success rate, mean quality rating, time-estimation precision, and count of prior assignments. Derived features: aggregate proficiency sum, years of professional experience, and number of overlapping skills. Categorical variables were label-encoded; continuous variables were standardized via StandardScaler; absent values were imputed using the median (numeric) or mode (categorical).

Conclusions and Future Work

This article introduces SmartAssign, an intelligent recommendation system that matches software development tasks to developers through a four-factor hybrid algorithm. The weights are distributed as follows: skill

alignment 40%, performance history 30%, workload balance 20%, and seniority match 10%. The random forest model achieved an accuracy of 82.7% and a 0.891 ROC-AUC on a 500-record benchmark. The systematic comparison of eight classifiers confirms the superiority of ensemble methods for this prediction task.

From a practical perspective, under AI guidance, the system increased the task completion rate from a baseline of 70% to 85%, kept developer utilization within a healthy range of 55–85%, and shortened the average allocation review time by 86%. The explanation module—which generates a plain-language rationale for each recommendation—was identified as the most important driver of management trust.

Planned extensions include: (1) piloting SmartAssign with industry partners and running controlled A/B experiments comparing the algorithm against manual allocation; (2) modeling team interaction dynamics and communication-graph features to capture interpersonal factors; (3) deploying an online learning mechanism that gradually updates the model as new outcome data accumulates; (4) building connectors to the Jira and GitHub APIs so that the tool can be seamlessly integrated into established workflows; and (5) investigating graph neural networks over a developer-task-skill knowledge graph to learn richer latent representations. Supporting multi-project environments—where developers contribute to several teams simultaneously—is another natural extension that addresses widespread organizational realities. We plan to release an open-source version to encourage community adoption and further academic research.

References:

1. Standish Group, "CHAOS Report 2020: Beyond Infinity," The Standish Group International, 2020.
2. N. Benabbou, M. Lemaître, and X. Lorca, "Solving project scheduling problems with multi-skill workforce using constraint programming," *Constraints*, vol. 23, no. 4, pp. 447–474, 2018.
3. W. Maalej, R. Tiarks, T. Roehm, and R. Koschke, "On the comprehension of program comprehension," *ACM Trans. Softw. Eng. Methodol.*, vol. 23, no. 4, pp. 1–37, 2014.
4. J. Bobadilla, F. Ortega, A. Hernando, and A. Gutiérrez, "Recommender systems survey," *Knowl.-Based Syst.*, vol. 46, pp. 109–132, 2013.
5. G. Linden, B. Smith, and J. York, "Amazon.com recommendations: Item-to-item collaborative filtering," *IEEE Internet Comput.*, vol. 7, no. 1, pp. 76–80, 2003.
6. J. Anvik, L. Hiew, and G. C. Murphy, "Who should fix this bug?" in *Proc. 28th Int. Conf. Softw. Eng.*, 2006, pp. 361–370.
7. G. Jeong, S. Kim, and T. Zimmermann, "Improving bug triage with bug tossing graphs," in *Proc. 7th Joint Meeting ESEC/FSE*, 2009, pp. 111–120.
8. A. Tamrawi, T. T. Nguyen, J. M. Al-Kofahi, and T. N. Nguyen, "Fuzzy set and cache-based approach for bug triaging," in *Proc. 19th ACM SIGSOFT Symp./13th Euro. Conf. FSE*, 2011, pp. 365–375.
9. W. Diehl and R. Souza, "Developer recommendation for issue solving in software projects: A systematic literature review," *Inf. Softw. Technol.*, vol. 114, pp. 1–18, 2019.
10. X. Xia, D. Lo, X. Wang, and B. Zhou, "Accurate developer recommendation for bug resolution," in *Proc. 20th WCRE*, 2013, pp. 72–81.
11. C. Molnar, *Interpretable Machine Learning: A Guide for Making Black Box Models Explainable*, 2020. Online. Available: <https://christophm.github.io/interpretable-ml-book/>
12. M. T. Ribeiro, S. Singh, and C. Guestrin, "Why should I trust you? Explaining the predictions of any classifier," in *Proc. 22nd ACM SIGKDD*, 2016, pp. 1135–1144.
13. F. Pedregosa et al., "Scikit-learn: Machine learning in Python," *J. Mach. Learn. Res.*, vol. 12, pp. 2825–2830, 2011.
14. T. Chen and C. Guestrin, "XGBoost: A scalable tree boosting system," in *Proc. 22nd ACM SIGKDD*, 2016, pp. 785–794.
15. G. Ke et al., "LightGBM: A highly efficient gradient boosting decision tree," in *Advances in NIPS*, vol. 30, 2017, pp. 3146–3154.
16. J. Marlow, L. Dabbish, and J. Herbsleb, "Impression formation in online peer production: Activity traces and personal profiles in GitHub," in *Proc. CSCW*, 2013, pp. 117–128.