



A MACHINE LEARNING-BASED DECISION-SUPPORT PLATFORM FOR SUITABILITY CLASSIFICATION OF BIODEGRADABLE WOUND-DRESSING MATERIALS

¹Mrs. P. Rajeshwari, ²A. G. Senbagavalli, ³Santhosh B, ⁴Sadhurmitha V P, ⁵Sasmitha R

¹Assistant Professor, Department of Software System and AIML, ^{2,3,4,5}Student, B.Sc AIML Sri Krishna Arts and Science College

Abstract: Choosing a biodegradable wound-dressing formulation that meets mechanical strength, controlled breakdown, sufficient porosity, a biocompatible pH, and antimicrobial effectiveness all at once is traditionally done through slow, resource-heavy lab testing across an enormous space of candidate materials and coatings. This paper introduces BEST (Biomaterial and Engineering Suitability Tool), an end-to-end decision-support platform that classifies a candidate wound-dressing formulation using ten physico-chemical input measurements, without needing the material's identity to be specified. The system pairs a soft-voting ensemble made up of a Random Forest model and a Gradient Boosting model, trained on a domain-informed synthetic dataset of 2,000 samples, inside one scikit-learn pipeline, and serves it through a Flask REST API consumed by a React single-page front end. In addition to a predicted class, the system produces a property-level rating, an explanation based on feature importance, and a templated natural-language suggestion for the next experimental step. The deployed ensemble reaches 88.0% accuracy on the held-out test set and 85.2% (plus or minus 1.1%) across five-fold cross-validation, with the antimicrobial index and biodegradation rate standing out as the strongest predictors. The platform additionally includes several decision-support add-ons — an explainable-AI feature-impact display, a molecular structure viewer for bioactive coating compounds, an automated formulation optimizer, a sustainability (BioScore) metric, a side-by-side formulation comparison tool, and an exportable PDF report — turning it into a full first-pass screening assistant rather than a plain classifier. This paper covers the motivation, dataset construction, system architecture, implementation, results, limitations, and directions for future research on the platform.

Index Terms - Machine Learning, Biomaterials, Wound Dressing, Ensemble Learning, Decision Support, Explainable AI

I. INTRODUCTION

Materials used in wound dressings need to meet a demanding set of mechanical, chemical, and biological requirements at the same time: enough tensile strength to survive handling, a biodegradation rate that lines up with the natural healing timeline, sufficient porosity and water absorption to handle exudate, a pH compatible with biological tissue, and antimicrobial ability that controls infection without causing cytotoxic effects. Finding formulations that satisfy all of these criteria together is traditionally done through repeated laboratory testing, a process that takes considerable time, consumes significant resources, and is hard to scale given the enormous number of possible material and coating combinations.

Machine learning (ML) has been increasingly adopted within materials science as a way to shorten this repeated testing cycle, by learning the relationship between measurable material properties and a target performance label, so that candidate formulations can be evaluated computationally before any physical

sample is made. Tree-based ensemble methods have shown particular promise on the small-to-medium, tabular, and frequently noisy datasets common in experimental materials science, because they blend the variance-reducing effect of bagging with the bias-reducing effect of boosting.

This paper describes the design, build, and evaluation of BEST (Biomaterial and Engineering Suitability Tool), an end-to-end platform that classifies a candidate wound-dressing formulation directly from ten physico-chemical descriptors, without needing to know what the material actually is. The key contributions of this work are:

- 1) A domain-informed methodology for generating a synthetic dataset that turns ten physico-chemical descriptors into a weighted composite suitability score using ideal ranges drawn from the literature.
- 2) A soft-voting ensemble combining a Random Forest and a Gradient Boosting model within a single scikit-learn pipeline, light enough to be retrained without specialized hardware.
- 3) A complete Flask-plus-React deployment that returns a property-by-property rating and a natural-language recommendation together with the prediction.
- 4) A collection of decision-support extensions — an explainable-AI feature-impact display, a molecular structure viewer, an automated optimizer, a sustainability BioScore, a comparison tool, and an exportable PDF report.
- 5) An analysis of feature importance together with a discussion of the limits of biomaterial classifiers trained on synthetic data, including directions for moving toward curated experimental data.

The rest of the paper proceeds as follows. Section II covers related work. Section III describes the dataset and the chosen features. Section IV presents the system's architecture and modelling approach. Section V discusses the implementation. Section VI reports the experimental results. Section VII discusses the platform's advantages, limitations, and how it compares with traditional screening approaches, and Section VIII closes with directions for future work.

II. RELATED WORK

A large body of prior work surveys the materials used in modern wound dressings, ranging from films, foams, hydrocolloids, alginates, and hydrogels [1] to multifunctional “smart” dressings that incorporate embedded sensors [3], [4] and biomimetic nano-scale dressings [5], [7]. Related reviews list the essential physico-chemical properties an effective dressing needs to have — biocompatibility, controlled breakdown, and adequate mechanical support — and point out that natural and synthetic polymers each balance antimicrobial activity, mechanical stability, and degradation control in different ways [6].

On the computational side, machine learning has been used in materials science both to predict how a material will perform based on its composition, and in the reverse direction, to generate or refine candidate designs aimed at a target property [2]. Sensor-based smart dressings have also been suggested as a way to continuously collect data for predicting wound healing, although these systems track an existing dressing that is already in place rather than screening candidate formulations before they are made [4]. The ensemble techniques used in this work rest on well-established foundations: Random Forests reduce variance by averaging decorrelated trees built on bootstrapped samples [8]; Gradient Boosting reduces bias on tabular data by growing an additive sequence of trees fit to the residual error [9]; both methods, along with the soft-voting classifier that combines them, are available in scikit-learn [10].

Unlike the work summarized above, BEST is aimed at an earlier, pre-experimental screening stage: given measured or estimated physico-chemical descriptors for a candidate formulation, it produces an interpretable suitability judgement along with guidance on which properties are limiting — a low-cost, first-pass decision-support step that comes before in-vitro validation, rather than a substitute for laboratory testing.

III. DATASET AND FEATURE DESIGN

Ten physico-chemical descriptors were chosen as inputs to the model because of their well-established relevance to wound-dressing performance in the biomaterials literature [1], [6]: tensile strength, biodegradation rate, coating concentration, pH level, porosity, water absorption, elongation at break, antimicrobial index, surface roughness, and thermal stability. Table I lists each feature along with the literature-informed ideal range used in the property-rating logic.

TABLE I
LITERATURE-INFORMED IDEAL RANGES FOR THE TEN INPUT DESCRIPTORS

Feature	Unit	Ideal Range
Tensile Strength	MPa	40 – 150
Biodegradation Rate	days	7 – 60
Coating Concentration	%	1 – 5
pH Level	—	6.5 – 7.5
Porosity	%	60 – 85
Water Absorption	%	50 – 200
Elongation at Break	%	10 – 50
Antimicrobial Index	/100	≥ 60
Surface Roughness	μm	0.5 – 3
Thermal Stability	$^{\circ}\text{C}$	≥ 150

Since no sufficiently large, publicly available, labelled dataset of wound-dressing formulations currently exists, a synthetic dataset of 2,000 samples was created instead. Each feature was drawn independently from a uniform distribution spanning a physically reasonable range, and a normalized sub-score between 0 and 1 was calculated for each one, peaking at the ideal value taken from the literature and falling off the further a value moves away from it (for the antimicrobial index and thermal stability, the sub-score instead rises monotonically rather than peaking). These ten sub-scores were then combined through a weighted linear sum, with weights chosen to reflect clinical importance, to produce a single composite suitability score, to which a small amount of Gaussian noise was added. Samples were then labelled Suitable when the score was 0.62 or above, Partially Suitable for scores between 0.40 and 0.62, and Not Suitable below 0.40.

Within the generated dataset, 1,432 samples (71.6%) received the Suitable label and 568 (28.4%) were labelled Partially Suitable; no sample fell into the Not Suitable range, so in practice the deployed classifier operates as a binary classifier across the two classes that actually occurred (Fig. 1) — a point examined further in Section VII. Using this synthetic-data approach makes it possible to train the classifier end-to-end while curated experimental data is still unavailable, and it keeps the scoring rationale transparent and traceable to the literature, which matters for a tool meant to support expert judgement rather than replace it.

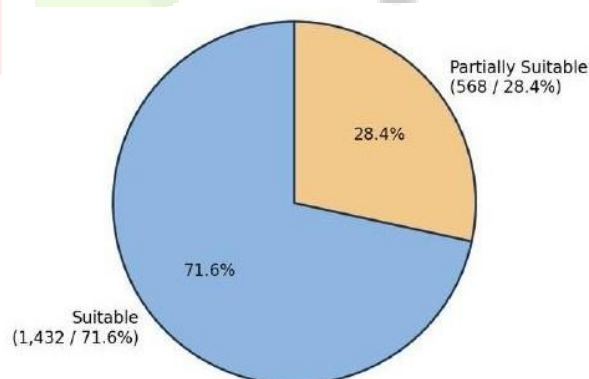


Fig. 1. Class distribution across the 2,000-sample generated dataset (71.6% Suitable / 28.4% Partially Suitable).

IV. SYSTEM ARCHITECTURE AND METHODOLOGY

A. Overall Architecture

BEST is built as a standard three-tier system: a React single-page client gathers the ten physico-chemical measurements along with free-text names for the material and coating; a Flask REST API checks the submitted values, converts them into numeric form, and passes them to the inference pipeline; and the inference pipeline itself — a serialized scikit-learn Pipeline — standardizes the features before handing them to the trained ensemble classifier. The predicted class, the class probabilities, a rule-based rating for each property, and a templated natural-language recommendation are all put together on the server side and returned as one JSON response.

B. Model Training Pipeline

The modelling process starts from the synthetic dataset described in Section III. The data is split 80/20 into training and test sets (1,600 training samples, 400 test samples) using stratification to keep class proportions consistent, and a StandardScaler then normalizes each feature to zero mean and unit variance. Two base models are trained side by side: a Random Forest made up of 300 trees, with a maximum depth of 15, a minimum of two samples required per leaf, and balanced class weighting; and a Gradient Boosting classifier with 200 estimators, a maximum depth of 6, and a learning rate of 0.1. These two models are then combined through a soft-voting classifier, which averages their predicted class probabilities instead of taking a simple majority vote, so that whichever model is more confident on a particular sample has proportionally more influence on the final decision. The full pipeline is assessed both on the held-out test split and through five-fold cross-validation, and is then serialized for deployment along with the fitted label encoder.

V. IMPLEMENTATION

A. Backend and Frontend

The backend runs on Flask and exposes three endpoints: /predict takes in the ten numeric features along with optional material and coating names, and returns the predicted label, the confidence for each class, a Good/Moderate/Poor rating for each property computed against the ranges given in Table I, and a templated text recommendation; /metadata exposes the model's stored accuracy figures, cross-validation statistics, and feature importances; and /health serves as a basic liveness check. The client side is a single-page React application, made up of a Header, an InputForm containing numeric fields for each of the ten features plus two free-text fields for the material and coating name, a ResultPanel that displays the predicted label, the confidence level, and the class probabilities, and a ModelInfo component that shows the model's reported accuracy and its most influential features.

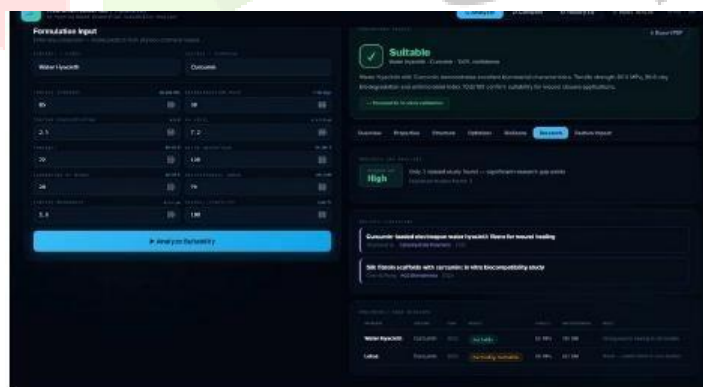


Fig. 2. The BEST web interface: formulation input on the left, prediction result with class probabilities and a property radar chart on the right.

B. Extended Decision-Support Modules

Beyond the core classification function, the platform includes several modules that expand a plain label into a richer decision-support experience. An explainable-AI feature-impact view shows, for each formulation that is submitted, which descriptors had the strongest effect on that specific prediction, drawing on the Random Forest's feature importances rather than only a single global ranking. A molecular structure viewer displays the structure of bioactive coating compounds — curcumin, for instance — so a user can check the chemical basis behind the modelled antimicrobial contribution. An optimizer module checks the formulation against the ideal ranges and flags any property that falls outside its acceptable band. A BioScore module estimates the relative environmental impact and carbon footprint of a formulation based on its material and coating makeup. A side-by-side comparison view lines up the property ratings, suitability classes, and BioScores for two or more candidate formulations, and the complete evaluation can be exported as a formatted PDF report.

VI. RESULTS AND DISCUSSION

Table III summarizes how the deployed voting ensemble performed. On the 400-sample held-out test split, the ensemble reached 88.0% accuracy; five-fold cross-validation across the full dataset gave a mean accuracy of 85.2% with a standard deviation of 1.1%, showing that performance stays stable across different train-test splits rather than being a result of one particularly favourable split.

TABLE III
HELD-OUT AND CROSS-VALIDATED CLASSIFICATION PERFORMANCE OF THE DEPLOYED ENSEMBLE

Metric	Value	Samples
Held-out Test Accuracy	88.0%	400
5-fold CV Accuracy	85.2% $\pm$ 1.1%	2,000
Training Set Size	—	1,600

Looking at the feature-importance results from the Random Forest component (Fig. 3), the antimicrobial index (0.201) and biodegradation rate (0.172) stand out as the two strongest predictors, with pH (0.108) and coating concentration (0.097) coming next. This lines up with clinical expectations: infection control and a degradation timeline that matches the healing process are widely seen as primary factors in dressing suitability within the biomaterials literature [1], [6], while secondary mechanical properties such as water absorption (0.044) and thermal stability (0.050) contribute comparatively less — likely because they less often act as the binding constraint within the ranges sampled to build the training labels.

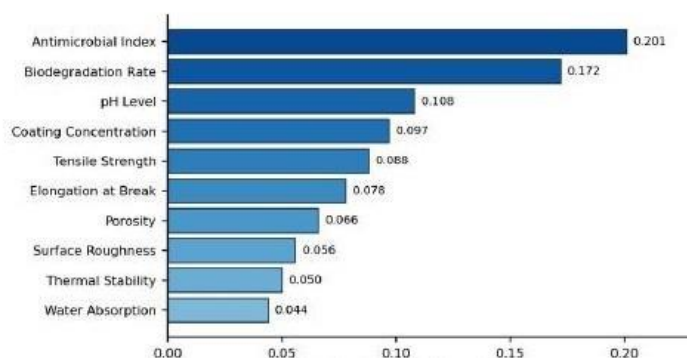


Fig. 3. Random Forest feature importances for the ten input descriptors.

Table IV shows an example prediction returned by the /predict endpoint for a formulation with near-ideal properties (tensile strength 85 MPa, biodegradation 30 days, pH 7.2, antimicrobial index 70/100). The system correctly labels the formulation as Suitable with 91.4% confidence, and each of the property-level ratings correctly comes back as Good, showing that the rule-based rating logic and the learned classifier agree with each other.

TABLE IV
REPRESENTATIVE /PREDICT RESPONSE FOR A NEAR-IDEAL CANDIDATE FORMULATION

Field	Value	Rating
Predicted Class	Suitable	—
Confidence	91.4%	—
Tensile Strength	85 MPa	Good
Antimicrobial Index	70/100	Good
Biodegradation Rate	30 days	Good
pH Level	7.2	Good

VII. DISCUSSION: BENEFITS, LIMITATIONS, AND COMPARISON

By giving an instant suitability judgement directly from descriptor values, the platform lets candidate formulations be triaged computationally before any physical fabrication or lab testing takes place, cutting down on wasted material and technician time. Rather than returning just a class label, it reports property-by-property ratings and a breakdown of feature impact, and the BioScore and carbon-footprint indicators build sustainability directly into the screening process. Because the ensemble is lightweight and works from tabular descriptors rather than raw imaging or spectroscopic data, the pipeline can be retrained and redeployed without any specialized hardware, which makes the tool usable by smaller labs; a suitability judgement that would take days to weeks in a physical lab is instead returned in seconds, with the same rule-based criteria applied consistently to every formulation submitted.

That said, several limitations need to be kept in mind when interpreting the reported accuracy. Most importantly, the training data is synthetically generated from a hand-designed weighted scoring rule rather than from measured biological outcomes, so the reported accuracy reflects how well the ensemble reproduces that scoring rule rather than genuine clinical or in-vitro suitability; replacing or supplementing the synthetic dataset with curated experimental data is a necessary next step. Second, because no sample fell below the Not Suitable threshold under the sampling approach used, the deployed model was effectively trained and tested as a two-class classifier even though the interface was built with three classes in mind — so the reported accuracy does not reflect how well the model recognizes genuinely unsuitable formulations, which is arguably the more safety-critical distinction. Third, the ten features were sampled independently of each other, so real-world correlations between properties (such as between porosity and water absorption) are not captured in the training distribution. Finally, as with any computational pre-screening tool, the outputs would need thorough in-vitro and eventual clinical validation before being used to inform actual formulation decisions; the platform is intended strictly as a first-pass, low-cost decision-support aid rather than a regulatory-grade diagnostic tool, and it cannot directly confirm biological outcomes such as cytotoxicity or in-vivo healing performance.

VIII. FUTURE WORK AND CONCLUSION

Future work will focus first on replacing the synthetic training data with a dataset built from published experimental characterizations of real wound-dressing materials, and on deliberately including boundary and clearly unsuitable formulations so the training distribution genuinely spans three classes. Applying model-agnostic explainability methods such as SHAP would add per-prediction attribution on top of the existing feature importances, and extending the task from classification to a regression estimate of expected healing or degradation time would provide more fine-grained decision support. The BioScore and carbon-footprint indicators could be extended using life-cycle-assessment data specific to individual biomaterials, and further ahead, adding imaging or spectroscopic characterization data alongside the current tabular descriptors [13] could broaden the range of evidence the platform is able to evaluate.

This paper presented BEST, a full-stack machine learning platform for automated, interpretable suitability screening of biodegradable wound-dressing formulations based on ten physico-chemical descriptors. A soft-voting ensemble combining a Random Forest and a Gradient Boosting classifier, trained on a domain-informed synthetic dataset, reaches 88.0% held-out accuracy and 85.2% ($\pm 1.1\%$) cross-validation accuracy, and is deployed behind a Flask REST API consumed by a React interface that returns a classification together with a property-by-property explanation, an explainable-AI feature-impact view, a molecular structure viewer, an automated optimizer, a sustainability BioScore, a comparison tool, and an exportable PDF report. The feature-importance analysis shows that antimicrobial performance and

biodegradation timing are the strongest predictors of suitability, consistent with established biomaterials literature. Future work will concentrate on replacing the synthetic dataset with curated experimental data and on strengthening the platform's explainability and sustainability-modelling components, moving the tool from an illustrative prototype toward a validated, laboratory-ready decision-support instrument.

REFERENCES

- [1] "Synthetic polymeric biomaterials for wound healing: a review," Progress in Biomaterials, Springer Nature, 2018.
- [2] "A review on the applications of machine learning in biomaterials, biomechanics, and biomanufacturing for tissue engineering," ScienceDirect, Elsevier, 2025.
- [3] "Multifunctional polymeric wound dressings," Polymer Bulletin, Springer Nature, 2025.
- [4] "Smart Biomaterials in Wound Healing: Advances, Challenges, and Future Directions in Intelligent Dressing Design," MDPI, 2025.
- [5] "Biomimetic nano dressing in wound healing: design strategies and application," Burns & Trauma, Oxford Academic, 2025.
- [6] "Natural and synthetic biomaterials, structural matrices-based wound dressings: Key properties, material correlation, and adaptability," ScienceDirect, Elsevier, 2025.
- [7] "Advances and Challenges in Immune-Modulatory Biomaterials for Wound Healing Applications," PubMed Central, National Institutes of Health.
- [8] L. Breiman, "Random Forests," Machine Learning, vol. 45, no. 1, pp. 5–32, 2001.
- [9] J. H. Friedman, "Greedy function approximation: A gradient boosting machine," Annals of Statistics, vol. 29, no. 5, pp. 1189–1232, 2001.
- [10] F. Pedregosa et al., "Scikit-learn: Machine Learning in Python," Journal of Machine Learning Research, vol. 12, pp. 2825–2830, 2011.
- [11] Pallets Projects, "Flask Documentation." [Online]. Available: <https://flask.palletsprojects.com>
- [12] Meta Open Source, "React Documentation." [Online]. Available: <https://react.dev>
- [13] G. Litjens et al., "A survey on deep learning in medical image analysis," Med. Image Anal., vol. 42, pp. 60–88, 2017.