

# Zdolność rozdzielcza gładkiego algorytmu ewolucyjnego

Marek Gutowski  
Instytut Fizyki PAN  
02-668 Warszawa, Al. Lotników 32/46  
e-mail: gutow@GAMMA1.ifpan.edu.pl

## STRESZCZENIE

Artykuł prezentuje wersję gładką algorytmu ewolucyjnego, opisywaną już wcześniej, oraz dodatkowo, wprowadza pojęcie *zdolności rozdzielczej*. Niniejsza praca wskazuje na konieczny związek między możliwą do uzyskania zdolnością rozdzielczą (stopniem szczegółowości rozwiązania) a licznnością danych wejściowych. Relacja ta wykazuje pewne podobieństwo do znanej z fizyki kwantowej zasady nieoznaczoności Heisenberga.

## WSTĘP

Pomysł opracowania algorytmów ewolucyjnych [1] wziął się z obserwacji Przyrody, która w zadziwiająco skuteczny sposób potrafi radzić sobie z problemem przystosowania do środowiska naturalnego. Na bazie podstawowych zasad rządzących ewolucją naturalną powstał schemat operujący na sztucznie stworzonych „osobnikach”. Tymi „osobnikami” (*chromosomami*) są zespoły specyficznie zakodowanych parametrów (*genów*), na których operuje algorytm ewolucyjny starając się je dobrać jak najlepiej do postawionego zadania. Tak jak czyni to Przyroda, algorytmy ewolucyjne poszukują optymalnego rozwiązania przy pomocy mechanizmów doboru naturalnego i dziedziczenia, wykorzystując przy tym darwinowską zasadę największych szans na przeżycie dla osobników najlepiej przystosowanych, oraz nieustanną, chociaż losową, wymianę informacji pomiędzy nimi.

Algorytmy ewolucyjne, zwane często genetycznymi, jako metody optymalizacyjne cieszą się wzrastającym powodzeniem ze względu na ich uniwersalność, skuteczność, elastyczność oraz łatwość implementacji w różnych językach programowania. Jak dotąd, najbardziej rozpowszechniona jest rodzina algorytmów genetycznych, która do opisu problemów używa zmiennych binarnych. Forma ta była rozszerzana przez wielu autorów (patrz np. [2]) tak, by możliwe było również operowanie niewiadomymi typu rzeczywistego (ciągłego).

Klasyczny algorytm genetyczny można zmodyfikować w taki sposób, zachowując jednocześnie wszystkie istotne cechy pierwowzoru, aby mógł on modelować krzywe ciągłe i *gładkie* [3]. Ta właściwość może być wykorzystana np. do filtracji (wygładzania) krzywych o nieznanym modelu analitycznym [4], albo do usuwania zbędnego tła w sygnale fizycznym.

## GLADKI ALGORYTM GENETYCZNY

Wersja gładka algorytmu znajduje naturalne zastosowanie w zagadnieniach, w których oczekuje się, że poszukiwane rozwiązanie powinno mieć postać gładkiej krzywej. Przykła-

dem może tu być np. praca [5], w której zadaniem algorytmu było znalezienie rozkładu wielkości ziaren magnetycznych (nanokryształów), którego charakterystyka magnesowania, tj. pętla histerezy magnetycznej w ustalonej temperaturze, byłaby identyczna z obserwowaną eksperymentalnie. Nie zakładano przy tym *a priori* żadnego konkretnego kształtu poszukiwanego rozkładu, pozwalając na zupełnie swobodną ewolucję początkowej populacji próbných rozwiązań.

Rozwiązanie problemu w postaci gładkiej krzywej oznacza w praktyce znalezienie uporządkowanego zbioru par liczb  $(x_i, y_i)$ ,  $i = 1 \dots n$ , które po zobrazowaniu w odpowiednim układzie współrzędnych tworzą gładką linię. Oczywiście jest to forma zdyskretyzowana pewnej gładkiej krzywej. Zwykle liczby  $x_i$  są zadane z góry lub przynajmniej znany jest przedział  $[x_{min}, x_{max}]$ , w którym powinny one się mieścić. W tym drugim przypadku pozostaje ustalić wartość liczby  $n$ , czyli ilość poszukiwanych niewiadomych, i stosownie — np. równomiernie — rozłożyć liczby  $x_i$  w tym przedziale. Tymczasem przyjmujemy, że liczby  $x_i$ , oraz ich ilość są znane i ustalone. Niewiadomymi, których naprawdę poszukujemy są liczby  $y_i$ , zaś liczby  $x_i$  służą jedynie do swoistego numerowania niewiadomych. Nie zakładamy przy tym równomiernego rozmieszczenia  $x_i$  na osi liczbowej — w wielu przypadkach bardziej przydatna może okazać się np. skala logarytmiczna. Powiedzmy od razu, że w problemach, w których poszukiwane niewiadome nie są ze sobą aż tak ściśle związane, albo kiedy ich kolejność może być wybrana dowolnie, (jak np. składowe rozwiązania wielowymiarowego układu równań) prezentowany algorytm będzie również działał, jeśli tylko położymy  $x_i \equiv i$ , co odpowiada zwyczajnemu ponumerowaniu niewiadomych.

Początkową populację osobników będących próbnymi rozwiązaniami tworzymy jako  $n$ -elementowe ciągi  $(y_1, y_2, \dots, y_n)$ ,  $(y'_1, y'_2, \dots, y'_n)$ , itd. Widać, że liczby  $y_i$  pełnią rolę pojedynczych genów. Populacji tej pozwolimy następnie swobodnie ewoluować z pokolenia na pokolenie, według klasycznych reguł, które teraz omówimy.

### GLADKIE OPERACJE GENETYCZNE

Jako pierwszą omówimy operację krzyżowania, czyli reprodukcji. Podobnie jak w wersji dyskretnej, dwa potencjalne osobniki rodzicielskie,  $r_1$  i  $r_2$ , albo zostają przepisane do nowego pokolenia bez żadnych zmian (przeżywają bezpotomnie), albo zostają zastąpione dwoma osobnikami potomnymi,  $p_1$  i  $p_2$ , otrzymywanymi jak następuje:

$$\begin{aligned} p_1(x_j) &= w(x_j) \cdot r_1(x_j) + [1 - w(x_j)] \cdot r_2(x_j) \\ p_2(x_j) &= [1 - w(x_j)] \cdot r_1(x_j) + w(x_j) \cdot r_2(x_j) \end{aligned} \quad (1)$$

Wprowadziliśmy w tych formułach funkcję  $w(x_j)$ , którą będziemy dalej nazywać funkcją mieszającą. Ma ona następujące właściwości:

1.  $w$  jest określona na przedziale  $[x_{min}, x_{max}]$ , tj. tym samym, w którym mieszczą się wszystkie geny  $y_i$ ,
2.  $0 \leq w(x_j) \leq 1$  dla  $1 \leq j \leq n$ ,
3.  $w$  jest gładka, lecz różna od stałej (stała funkcja  $w$  sprowadza algorytm do swoistej odmiany dobrze znanej metody sympleksów).

Wybieramy  $w$  w postaci:

$$w(x) = \frac{1}{2} \left[ 1 + \tanh \frac{x - \xi}{\Delta} \right] \quad (2)$$

Liczba  $\xi$  w powyższym wzorze jest dowolną liczbą z przedziału  $[x_{min}, x_{max}]$ , dobiebraną losowo, oddzielnie dla każdej pary osobników rodzicielskich. Jest to tzw. punkt krzyżowania ( $w(x = \xi) = \frac{1}{2}$ ). Tak dobrana funkcja  $w$  ma kształt rozmytego skoku jednostkowego zlokalizowanego w  $x = \xi$ . Wielkość rozmycia scharakteryzowana jest parametrem  $\Delta$ , równym – jak nietrudno się przekonać – różnicy odciętych punktów przecięcia stycznej do wykresu  $y = w(x)$  w punkcie (przebiegu)  $x = \xi$  a asymptotami funkcji  $w$ , tj. prostymi  $y = 0$  i  $y = 1$ . Wielkość  $\Delta$  nazywamy, przez analogię z zagadnieniami spektroskopowymi, **zdolnością rozdzielczą**.

O ile sposób krzyżowania chromosomów gładkich nie różni się drastycznie od przypadku zmiennych typu binarnego (tam funkcja mieszająca to po prostu skok jednostkowy) to procedura kreowania mutantów jest jakościowo różna. W przypadku gładkim nie możemy ograniczać się do modyfikacji pojedynczego genu – jednocześnie z nim modyfikacji muszą ulec również jego najbliżsi sąsiedzi. Operacja mutacji przebiega następująco:

$$o'(x_j) = o(x_j) \cdot \left[ 1 \pm \frac{1}{2} g(x_j - \xi) \right] \quad \text{dla } j = 1 \dots n \quad (3)$$

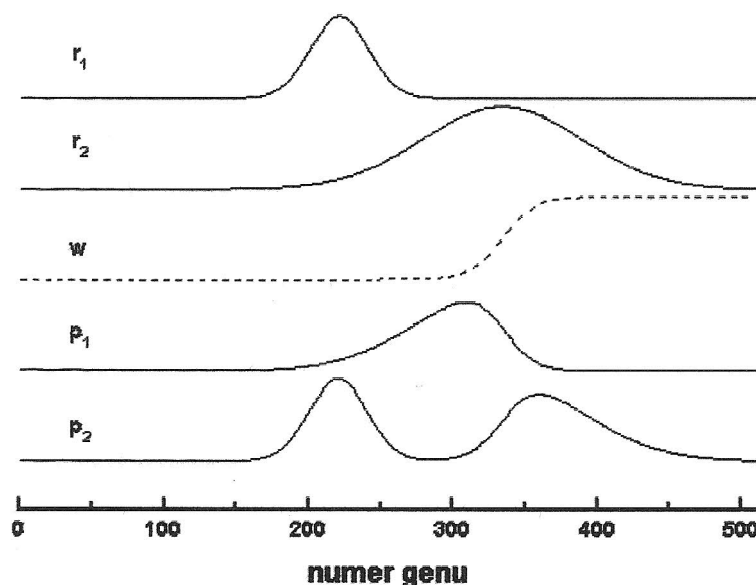
gdzie  $o()$  i  $o'()$  to odpowiednio osobnik przed i po mutacji, natomiast nowa funkcja  $g()$  ma tę własność, że różni się wyraźnie od zera jedynie w okolicy  $x = \xi$ , ponadto jest gładka, nieujemna i ma maksymalną wartość równą jedności. Warunki te spełnia np. odpowiednio znormalizowana funkcja Gaussa o niezbyt małej szerokości połówkowej. Znak w równaniu 3 dobierany jest losowo, podobnie jak liczba  $\xi$  – położenie (numer) najsilniej modyfikowanego genu (jest to oczywiście inna liczba losowa niż ta występująca w równaniu 2). Operacja mutacji zrealizowana w takiej postaci powoduje istotną modyfikację jedynie części chromosomu pozostawiając resztę genów praktycznie nie zmienioną. Szerokość połówkowa mutacji powinna wynosić co najmniej  $\Delta$  (patrz równanie 2). Zmianie musi też ulec określenie prawdopodobieństwa mutacji: zamiast liczby rzędu  $10^{-3}/\text{gen/pokolenie}$  wygodniej jest się posługiwać prawdopodobieństwem mutacji *całego* osobnika. Rozsądną wartością wydaje się taka, że w każdym pokoleniu zostaje zmutowany średnio jeden osobnik z całej populacji. Nieujemność funkcji  $g()$  jest postulowana jedynie dla zapewnienia uniwersalności algorytmu. Zapewnia to, że jeśli poszukiwana krzywa, reprezentująca np. rozkład prawdopodobieństwa, musi zachowywać stały znak, to własność ta zostaje utrzymana nawet po tak drastycznej operacji jak mutacja.

Zauważmy, że opisane operacje krzyżowania i mutacji dają gwarancję, że potomstwo pochodzące od gładkich rodziców będzie również gładkie.

## DOBÓR OSOBNIKÓW RODZIELSKICH

Zasada doboru potencjalnych osobników rodzicielskich jest taka sama jak zawsze: większe szanse mają osobniki lepiej przystosowane. To niezbyt precyzyjne określenie wyodrębnia z całej populacji pewien zbiór rozmyty, którego funkcja przynależności może być utożsamiona z prawdopodobieństwem zostania wylosowanym do procesu krzyżowania. Funkcja przynależności do tego zbioru musi oczywiście zależeć jednoznacznie od bieżącej wartości funkcji celu rozpatrywanego osobnika. Chętnie używamy następującego przekształcenia funkcji celu (przystosowania) na funkcję przynależności do *zbioru osobników dobrze przystosowanych*, podanego po raz pierwszy w [3]:

$$p_{doboru} = \frac{1}{1 + \exp \frac{f_{celu} - M_f}{S_f}} \quad (4)$$



Rys. 1: Ilustracja krzyżowania:  $r_1$  i  $r_2$  oznaczają rodziców,  $p_1$  i  $p_2$  – osobniki potomne, natomiast  $w$  jest kształtem funkcji mieszania (ze zdolnością rozdzielczą  $\Delta \approx 70$ ). Punkt krzyżowania odpowiada miejscu, w którym funkcja mieszania ma największe nachylenie. Widać dziedziczenie cech rodziców przez osobniki potomne.

We wzorze tym  $M_f$  oznacza medianę funkcji celu wszystkich osobników z populacji, a  $S_f$  – dyspersję wartości funkcji celu wokół mediany. Wzór 4 ma zastosowanie, gdy poszukiwanym optimum jest *minimum* funkcji celu; przy poszukiwaniu *maksimum* należy zmienić znak argumentu funkcji exp.

Określenie przystosowania w powyższy sposób przypomina — nie przypadkowo — funkcję Fermiego-Diraca znaną z fizyki statystycznej. Wybór ten ma wiele zalet. Po pierwsze, zminimalizowany jest wpływ osobników wyjątkowo dobrze/źle przystosowanych na przebieg ewolucji, tym samym problem przedwczesnej zbieżności, który został wielokrotnie opisany, m.in. przez Goldberga [3], praktycznie przestaje istnieć. Z drugiej strony, następuje wyjątkowo skuteczne rozdzielenie osobników średnio przystosowanych na dwie grupy o różnych perspektywach dalszego powielania ich materiału genetycznego. W początkowych pokoleniach duża wartość  $S_f$  powoduje, że wszystkie osobniki mają prawie jednakowe szanse przeżycia, podczas gdy w końcowych stadiach optymalizacji  $S_f$  jest małe i tylko „lepsz połowa” osobników ma realne szanse przekazania swojego materiału genetycznego do następnego pokolenia. Jednocześnie, proces optymalizacji w żadnej swojej fazie nie zostanie zdominowany przez wyjątkowo dobrze/źle przystosowanego mutantą, gdyż nawet bardzo odbiegająca od przeciętnej wartość funkcji celu nie wpływa na wartość mediany  $M_f$ , manifestując się jedynie umiarkowanym zwiększeniem wartości parametru  $S_f$ .

Nie każda para potencjalnych rodziców dochowuje się potomstwa. Decyduje o tym zewnętrzny parametr  $p_r$ , zadawany przez użytkownika i określający prawdopodobieństwo zdarzenia polegającego na tym, że wylosowana para zostanie zastąpiona w kolejnym po-

koleniu przez swoje potomstwo. Trudno podać uniwersalną wartość  $p_r$ , optymalną dla dowolnie wskazanego problemu. Doświadczenia numeryczne wskazują, że  $p_r$  zawarte w granicach  $[0.25, 0.95]$  nie ma zauważalnego wpływu na przebieg procesu ewolucji ani na otrzymywane wyniki. Mniejsze wartości  $p_r$  czasami prowadzą do pogorszenia stabilności przebiegu ewolucji, tzn. z jednej strony ilość pokoleń potrzebnych do otrzymania rozwiązania wyraźnie wzrasta (co nie jest jednakże jednoznaczne ze wzrostem ilości obliczeń), z drugiej zaś daje się zaobserwować brak monotoniczności przebiegu *średniej* wartości funkcji przystosowania w czasie działania algorytmu. Jako „bezpieczną” wartość autor przyjmuje zwykle  $p_r = \frac{2}{3}$ .

Kryteria zatrzymania algorytmu genetycznego w wersji gładkiej są identyczne jak w innych jego odmianach, nie będziemy ich więc tu dyskutować.

### ZNACZENIE ZDOLNOŚCI ROZDZIELCZEJ

Posłużymy się zagadnieniem wziętym z pracy [6]. Zadanie polega na znalezieniu radialnego rozkładu gęstości prądu płynącego w nadprzewodzącej próbce o kształcie płaskiego pierścienia, umieszczonej w znanym zewnętrznym polu magnetycznym, zorientowanym prostopadłe do powierzchni pierścienia. Dane wejściowe to radialny rozkład pola magnetycznego (a ściślej jedynie jego składowej prostopadłej) przy powierzchni próbki. Pole to, wyznaczone eksperymentalnie w  $N$  punktach pomiarowych, jest wypadkową zadanego pola zewnętrznego oraz tego wytwarzanego przez prądy płynące w badanej próbce.

Przekład problemu na język algorytmów ewolucyjnych wygląda następująco:

- zakładamy, że poszukiwana gęstość prądu  $j(r)$  zależy jedynie od odległości od środka pierścienia,  $r \in [r_{min}, r_{max}]$ .
- zakładamy pożądaną dokładność (stopień szczegółowości) rozwiązania, tj. ilość liczb,  $n$ , z których składać się będzie poszukiwane rozwiązanie.
- ciągi o postaci  $[j(r_1 = r_{min}), j(r_2), \dots, j(r_n = r_{max})]$  będą odtąd używane jako próbne rozwiązania problemu (chromosomy). Liczby  $r_1, r_2, \dots, r_n$ , określające odległość od środka pierścienia jednocześnie numerują poszczególne geny.
- funkcję celu określamy np. jako sumę kwadratów różnic pomiędzy obserwowanymi wartościami pola magnetycznego a tymi, które można wyliczyć z założonego rozkładu prądów (wyliczenia/symulacje, w oparciu o prawo Biot-Savarta, prowadzone są dla tych samych punktów przestrzeni, w których uzyskano wartości eksperymentalne, które *nie muszą się pokrywać* z położeniami  $r_1, r_2, \dots$  wypisanymi wyżej, tym bardziej, że prądy nie mogą płynąć poza pierścieniem, podczas gdy pole magnetyczne można zmierzyć — i mierzy się — również w innych miejscach).

Widać pewną dowolność w doborze liczby  $n$ . Nietrudno się przekonać, np. uruchamiając program, że algorytm ewolucyjny jest w stanie wygenerować *jakiś* rozwiązania dla rozmaitych wartości  $n$ . Chciałoby się mieć rozwiązania możliwie „dokładne”, co skłania do używania jak największych  $n$ , ograniczonych jedynie dostępnymi zasobami komputera. Powstaje jednak intrygujące pytanie: jak to możliwe, że w wyniku procesu optymalizacyjnego otrzymujemy więcej informacji, niż zawierały je dane wejściowe?

Sprzeczność jest, na szczęście, pozorna. Efektywnie musi być, oczywiście,  $n < N$ . Najlepsza zdolność rozdzielcza, jaką możemy uzyskać to



$$N \cdot \Delta_{opt} \approx r_{max} - r_{min} \quad (5)$$

gdzie  $N$  oznacza ilość punktów pomiarowych, czyli danych wejściowych.

W konkretnych rachunkach można przyjąć inną wielkość zdolności rozdzielczej  $\Delta$ , pod warunkiem, że spełniona będzie nierówność  $\Delta > \Delta_{opt}$ . W zagadnieniach takich jak np. w pracy [7] będziemy zwykle zainteresowani, aby  $\Delta \gg \Delta_{opt}$ , gdyż jedynie wtedy skutecznie uda się wydzielić interesujący sygnał, będący serią krótkich impulsów, o amplitudzie charakteryzującej się stałym znakiem, chaotycznie rozmieszczonych na wolnozmiennym tle obciążonym „zwykłym” szumem. We wszystkich jednak przypadkach powinniśmy prowadzić proces optymalizacji z zachowaniem warunku

$$n \cdot \Delta \geq N \cdot \Delta_{opt}, \quad (6)$$

przy czym  $\Delta \geq \Delta_{opt}$  też jest ograniczone od dołu. Nadmierne zwiększanie liczby  $n$  jest niecelowe, gdyż nie daje *żadnych* dodatkowych informacji – można uważać, że tylko niektóre uzyskane liczby  $y(r_i)$  zawierają prawdziwą informację, natomiast pozostałe są jedynie wynikiem interpolacji. (W rzeczywistości nie da się wskazać, które  $y(r_i)$  są „godne zaufania”, a które jedynie interpolowane.)

Czym grozi niespełnienie nierówności 6? Zwykle tym, że otrzymane rozwiązanie będzie w całości, lub jedynie w części, „pomarszczone” („pofalowane”), podobnie jak to się dzieje przy aproksymacji danych pomiarowych wielomianem zbyt wysokiego stopnia (oscylacje). Innymi objawami użycia zaniżonej wartości  $\Delta$  mogą też być rozmaite artefakty pochodzące od przypadkowych niekorzystnych mutacji. Ich obecność, poważnie potraktowana, może prowadzić do całkowicie błędnych interpretacji otrzymanych wyników. Sam proces ewolucyjny zwykle trwa też wyraźnie dłużej, niż w warunkach optymalnych.

Na zakończenie jedna istotna uwaga. Pojęcie zdolności rozdzielczej, jeżeli ma zastosowanie do konkretnego problemu, *musi* być stosowane z całą konsekwencją, począwszy od kreowania początkowej populacji rozwiązań próbnych (chromosomów). Jeśli osobniki z pierwszego pokolenia będą zbyt urozmaicone, to proces ewolucyjny może zachować tę ich cechę, prowadząc do niepoprawnych rozwiązań w taki sam sposób jak źle zaaplikowane mutacje (tj. zmieniające w istotny sposób zbyt małą ilość genów). Trudno tu określić precyzyjnie, jakie chromosomy należy uważać za „zbyt urozmaicone”. Roboczo można przyjąć, że są to takie chromosomy, w których najmniejsza odległość pomiędzy sąsiednimi ekstremami, bądź punktami przegięcia, jest mniejsza niż  $\Delta_{opt}$ . Nie jest to jakieś szczególnie wyrafinowane i trudne do spełnienia żądanie, gdyż jak pokazano w pracy [4], dwie początkowe populacje chromosomów:

- złożona z funkcji stałych oraz
- złożona z linii prostych o różnych nachyleniach

prowadzą do bardzo zbliżonych i zadowalających rezultatów końcowych.

## WNIOSKI

Prezentowany algorytm jest intuicyjnie prosty, a mimo to niezwykle elastyczny i skuteczny. Z powodzeniem może być używany do problemów o bardzo różnym charakterze, zwłaszcza szeroko rozumianych spektralnych i spektroskopowych, choć nie tylko takich. Szczególnie przydatny będzie on wszędzie tam, gdzie precyzyjne wartości poszukiwanych parametrów są mniej istotne niż *kształt* poszukiwanej krzywej. Wprowadzone w

pracy pojęcie *zdolności rozdzielczej* ma zastosowanie przede wszystkim do problemów, w których osoba prowadząca obliczenia musi samodzielnie określić ilość poszukiwanych niewiadomych, choć nie jest ograniczone wyłącznie do takich zagadnień. W szczególności, pojęcie zdolności rozdzielczej będzie bardzo przydatne przy wygładzaniu danych eksperymentalnych dając, jako efekt uboczny, niekonwencjonalne możliwości detekcji i eliminacji tzw. *outliers*, czyli błędów grubych. Wygładzane dane nie muszą przy tym być typu „równoodległego” (wtedy lepsze mogą okazać się metody oparte na szybkiej transformacie Fouriera, FFT). Wprowadzenie zdolności rozdzielczej daje możliwość świadomej kontroli procesu optymalizacyjnego, którego pożądanym wynikiem jest możliwie dokładnie odtworzona poszukiwana krzywa, nie zawierająca jednak artefaktów.

Gładki algorytm ewolucyjny jest uniwersalnym narzędziem optymalizacji, jako że z równym powodzeniem może on być użyty do każdego zagadnienia optymalizacyjnego, w którym niewiadome są liczbami rzeczywistymi, tak jak np. w przypadku rozwiązywania układu równań. Jego zapotrzebowanie na pamięć operacyjną komputera jest rzędu  $\mathcal{O}(n)$ , gdzie  $n$  jest ilością niewiadomych, podczas gdy inne algorytmy wymagają nawet do  $\mathcal{O}(n^3)$  miejsc w pamięci. Oznacza to, że algorytm „lubi” problemy o dużym wymiarze (np. w pracy [7]  $n$  przekraczało 1400, donoszono też o rozwiązanych problemach z  $n = 8192$ ), choć z drugiej strony dokładnie ta sama procedura dała bardzo dobre wyniki w zagadnieniu opisanym zaledwie trzema parametrami. Mocną stroną algorytmu jest to, że operuje on wyłącznie wartościami funkcji celu, a nie np. jej pochodnymi. Z tego powodu wystarczy, aby funkcja celu była ciągła; przy czym nie jest to warunek konieczny.

## BIBLIOGRAFIA

- [1] J.H. Holland, *Adaptation in Natural and Artificial Systems*, The University of Michigan Press: Ann Arbor MI, 1975  
także: John H. Holland, *Algorytmy genetyczne*, Świat Nauki, wrzesień 1992, str. 34
- [2] David E. Goldberg, *Algorytmy genetyczne i ich zastosowania*, Wyd. Naukowo-Techniczne, Warszawa, 1995 (tłumaczenie z: *Genetic Algorithms in Search, Optimization and Machine Learning*, Addison-Wesley Publ. Co., Inc., Reading, Massachusetts, USA, 1989).
- [3] M.W. Gutowski, *Smooth genetic algorithm*, J. Phys. A: Math. Gen., **27**, 7893, 1994
- [4] Marcin Małys, *Zastosowanie algorytmów genetycznych do wygładzania danych eksperymentalnych o nieznanym modelu analitycznym*, praca magisterska, Politechnika Warszawska, Wydział FTiMS, Warszawa 1996
- [5] A. Ślawska-Waniewska, M. Gutowski and H.K. Lachowicz, *“Superparamagnetism in a nanocrystalline Fe-based metallic glass”*, Physical Review, **B46**, 14594, 1992
- [6] K. Piotrowski, M.U. Gutowska, M. Gutowski, *Magnetic flux density and radial distribution of currents in HTSC samples of specific shape determined by genetic algorithm*, International Conference on Substrate Crystals and HTSC Films, ICSC-F’96, September 16–20, 1996, Jaszowiec, Poland, poster PII-18, pełny tekst przyjęty do druku w Acta Physica Polonica
- [7] Marek Gutowski, Marcin Małys, *Gładki algorytm genetyczny wydziela szumy Barkhausena z danych pomiarowych*, III Krajowa Konferencja „Komputerowe wspomaganie badań naukowych”, Polanica Zdrój, 17–19 października 1996, materiały str. 125